

J. H. Hise

MONATSHEFTE
FÜR
MATHEMATIK
UND
PHYSIK

MIT UNTERSTÜTZUNG DES
ÖSTERR. BUNDESMINISTERIUMS FÜR UNTERRICHT
HERAUSGEGEBEN
VON
H. HAHN UND W. WIRTINGER
IN WIEN

XXXVI. BAND
1. HEFT

LEIPZIG 1929
AKADEMISCHE VERLAGSGESELLSCHAFT M. B. H.

*Alle für die Redaktion bestimmten Mitteilungen
und Sendungen sind zu adressieren:*

*Redaktion der Monatshefte für Mathematik
und Physik, Wien, IX., Strudlhofgasse Nr. 4
Mathematisches Seminar der Universität*

Inhaltsverzeichnis.

	Seite
Ayres W. L., On continua which are disconnected by the omission of any point and some related problems	135
Brouwer L. E. J., Mathematik, Wissenschaft und Sprache	153
Burstin C., Beiträge zur mehrdimensionalen Differentialgeometrie	97
Duschek A., Die Starrheit der Eiflächen	131
Furtwängler Ph., Über affektfreie Gleichungen	89
Haack W., Affine Differentialgeometrie der Strahlensysteme	47
Künne H., Ein Theorem der Kurventheorie	149
Kuratowski C., Théorème sur trois continus	77
Levi-Civita T., Über dynamische Beanspruchung elastischer Systeme	165
Mayer W., Über abstrakte Topologie	1
Moore R. L., Concerning upper semi-continuous collections	81
Schuntner E., Über eine Verallgemeinerung des Poissonschen Theorems	43

Literaturberichte.]

	Seite
Bauer-v. Hanxleden-Krause, Lehrbuch der Mathematik. Dintzl	29
Baldus R., Nichteuklidische Geometrie. Helly	23
Baur F., Korrelationsrechnung. Schrutka	25
Beck, Einführung in die Axiomatik der Algebra. Rella	18
Beck H., Elementargeometrie. Hofreiter	24
Becker O., Mathematische Existenz. Hahn	8
Bieberbach-Bauer, Vorlesungen über Algebra. Rella	18
Bieberbach L., Differential- und Integralrechnung. Hahn	11
Bommersheim P., Beiträge zur Lehre von Ding und Gesetz. Carnap	28
Briefwechsel zwischen Carl Friedrich Gauss und Christian Ludwig Gerling Wirtinger	3
Carathéodory C., Vorlesungen über reelle Funktionen. Hahn	10
Die Kegelschnitte des Apollonios. Wirtinger	3
Dickson L. E., Algebren und ihre Zahlentheorie. Rella	21
Eisenhardt L. P., Non-Riemannian Geometrie. Mayer	24
Enzyklopädie der mathematischen Wissenschaften. Mayer	25
Euler L., Drei Abhandlungen über die Auflösung der Gleichungen. Hofreiter	11
Fettweis E., Das Rechnen der Naturvölker. Wirtinger	3
Fraenkel A., Einleitung in die Mengenlehre. Menger	5
Gauss C. F., Bestimmung der Anziehung eines elliptischen Ringes. Wirtinger	11
Grafmann H., Projektive Geometrie der Ebene. Knoll	22
Greinacher H., Ionen und Elektronen. Sexl	26
Haas A., Materiewellen und Quantenmechanik. Sexl	25
Haas A., Die Welt der Atome. Sexl	25

Über abstrakte Topologie.

Von Walther Mayer in Wien.

Diese Arbeit heißt abstrakte Topologie, weil die Komplexe, von deren Eigenschaften gesprochen wird, keineswegs irgend welche geometrische Gebilde sein müssen, sondern abstrakte Dinge sind, die einem sechs Axiome umfassenden Systeme zu genügen haben.

Wenn also auf Anschaulichkeit auch Verzicht geleistet wurde, so gewinnen dadurch die Überlegungen und Resultate an Exaktheit.

Der erste Abschnitt definiert die wichtigsten Invarianten eines Komplexringes, wobei ich einen Komplex als Ring bezeichne, sobald ich ihn als Inbegriff seiner Ecken, ein- und mehrdimensionalen Seiten auffasse.

Im zweiten Abschnitt wird dann der Begriff eines Unterringes eines gegebenen Ringes erklärt und anschließend der Komplementring und Randring eines solchen Unterringes definiert. Diese Begriffe gestatten dann geometrische Probleme zu behandeln, was im dritten und vierten Abschnitte geschieht.

Im dritten Abschnitt wird das Problem der Unterteilung eines gegebenen kombinatorischen Komplexringes behandelt, und zwar das inverse Problem, unter welchen Umständen man gewisse Komplexe eines gegebenen Komplexringes zu einem Komplex vereinigen darf, ohne die Homologiegruppen des Ringes zu verändern.

Der vierte Abschnitt behandelt einen Zusammensetzungssatz, nämlich die Frage nach den Homologiegruppen eines Ringes Σ , der als Summe zweier seiner Unterringe Σ_1 und Σ_2 gegeben ist. Von Interesse dürfte die Beziehung (96) zwischen den Bettischen Zahlen dieser Ringe sein.

Im zweiten Abschnitt wird ein Beweis einer Verallgemeinerung eines bekannten Poincaréschen Resultates skizziert.

In die Topologie wurde ich durch meinen Kollegen Vietoris eingeführt, dessen Vorlesung 1926/27 ich an der hiesigen Universität besuchte. In den vielen Gesprächen über dieses Gebiet hat mir Herr Vietoris eine Fülle Anregungen gegeben, für die ich ihm meinen besten Dank ausdrücke.

I. ABSCHNITT.

Der Komplexring Σ .

Wir definieren den Komplexring als einen Inbegriff von Dingen, die wir Komplexe nennen.

Jedem Komplex ist eine ganze Zahl ρ zwischen Null und n

$$0 \leq \rho \leq n$$

zugeordnet, die wir die Dimension des Komplexes nennen.

Die Zahl n heißt die Dimension des Ringes.

$K^{(\rho)}$ bezeichne einen ρ dimensionalen Komplex.

I. Axiom: Es bestehe eine Operation zwischen den Komplexen des Ringes Σ . Die ρ dimens. Komplexe in Σ ($0 \leq \rho \leq n$) seien in bezug auf jene Operation Elemente einer kommutativen Gruppe, $\{K^{(\rho)}\}$ bezeichnet.

Wir wollen als Operationszeichen das $+$ -Zeichen verwenden und demgemäß die Operation Addition nennen.

Das Nullelement der ρ dimens. Komplexe sei $0^{(\rho)}$ bezeichnet.

II. Axiom: In der Gruppe $\{K^{(\rho)}\}$ gebe es kein von $0^{(\rho)}$ verschiedenes Element endlicher Ordnung.

Aus $t k^{(\rho)} = 0^{(\rho)}$, t ganze Zahl $\neq 0$, folge also stets $k^{(\rho)} = 0^{(\rho)}$. Seien dann endlich viele, z. B. τ (τ ganze Zahl), ρ dimens. Komplexe $a_1^{(\rho)}, a_2^{(\rho)}, \dots, a_\tau^{(\rho)}$ von Σ gegeben, so bildet die Gesamtheit der Komplexe

$$p_i a_i^{(\rho)}, \quad i = 1, \dots, \tau,$$

p_i ganze (positive oder negative) Zahl eine Untergruppe von $\{K^{(\rho)}\}$.

III. Axiom: Es gebe in Σ zu jedem ρ ($0 \leq \rho \leq n$) ein System von endlich vielen ρ dimens. Komplexen

$$(1) \quad a_1^{(\rho)}, a_2^{(\rho)}, \dots, a_\tau^{(\rho)}$$

der Eigenschaft, daß jeder ρ dimens. Komplex in Σ in der Gesamtheit

$$p_i a_i^{(\rho)}, \quad i = 1, \dots, \tau$$

enthalten sei.

Unter allen Systemen (1) dieser Eigenschaft gibt es solche mit einer Minimalzahl von Elementen. Sei α_ρ diese Zahl (sie ist bei gegebenem Ring Σ eine Funktion von ρ), so wollen wir ein solches System

$$(1') \quad a_1^{(\rho)}, a_2^{(\rho)}, \dots, a_{\alpha_\rho}^{(\rho)}$$

eine Basis der ρ dimens. Komplexe von Σ nennen.

¹⁾ Über gleiche Indizes ist — wo nicht ausdrücklich das Gegenteil verlangt wird — stets zu summieren.

Wir zeigen nun, daß jeder ρ dimens. Komplex $K^{(e)}$ in Σ in eindeutiger Weise mittels der Basis dargestellt wird.

Aus $K^{(e)} = p_i a_i^{(e)} = q_i a_i^{(e)}$, $i = 1, \dots, \alpha_e$, folgt $0^{(e)} = (p_i - q_i) a_i^{(e)} = \pi_i a_i^{(e)}$, wo $\pi_i = p_i - q_i$ bedeutet. Nun folgt aber aus $0^{(e)} = \pi_i a_i^{(e)}$, $\pi_i = 0$. Denn durch unimodulare Transformation der $a_i^{(e)}$ können wir $0^{(e)} = \pi_i a_i^{(e)}$ in die kanonische Form $E_1 \bar{a}_1^{(e)} = 0^{(e)}$ bringen, wobei $\bar{a}_i^{(e)}$, $i = 1, \dots, \alpha_e$ eine neue Basis der ρ dimens. Komplexe in Σ und E_1 der gemeinsame Teiler der Zahlen π_i ist. Aus $E_1 \bar{a}_1^{(e)} = 0^{(e)}$ folgt aber nach Axiom II entweder $\bar{a}_1^{(e)} = 0^{(e)}$, was unmöglich ist, da es sonst eine Basis der ρ dimens. Komplexe in Σ von $\alpha_e - 1$ Elementen gäbe ($\bar{a}_j^{(e)}$, $i = 2, \dots, \alpha_e$), oder $E_1 = 0$, d. h. dann $\pi_i = 0$, q. e. d.

Die Eigenschaft einer Basis, die Komplexe eindeutig darzustellen, ist für die Basissysteme charakteristisch.

Denn sei $a_1^{(e)}, \dots, a_t^{(e)}$ ein System dieser Eigenschaft, so folgt aus $0^{(e)} = \pi_i a_i^{(e)}$, $i = 1, \dots, t$ jedenfalls $\pi_i = 0$, $i = 1, \dots, t$.

Wäre nun dieses System keine Basis, $t > \alpha_e$, und $\bar{a}_1^{(e)}, \dots, \bar{a}_{\alpha_e}^{(e)}$ eine Basis der ρ dimens. Komplexe, so gilt jedenfalls $a_i^{(e)} = \sigma_{ij} \bar{a}_j^{(e)}$, $i = 1 \dots t$, $j = 1 \dots \alpha_e$.

Ferner ist $\pi_i a_i^{(e)} = \pi_i \sigma_{ij} \bar{a}_j^{(e)} = \bar{\pi}_j \bar{a}_j^{(e)}$ und es wäre bereits $0^{(e)} = \pi_i a_i^{(e)}$, wenn $\bar{\pi}_j = \pi_i \sigma_{ij} = 0$, $i = 1 \dots t$, $j = 1 \dots \alpha_e$ ist. Wegen $t > \alpha_e$ gibt es aber immer Lösungen von $\pi_i \sigma_{ij} = 0$ mit $\pi_i \neq 0$, d. h. aus $0^{(e)} = \pi_i a_i^{(e)}$ folgte dann nicht $\pi_i = 0$. (Widerspruch.)

Somit gilt der

Satz I: Durch ein Basissystem (1') ist jeder ρ dimens. Komplex $K^{(e)}$ eindeutig darstellbar; und ist jeder Komplex $K^{(e)}$ durch ein System (1') eindeutig darstellbar, so ist (1') ein Basissystem.

Ferner gilt Satz II: Zwischen zwei Basissystemen der ρ dimens. Komplexe

$$a_i^{(e)} \text{ und } \bar{a}_i^{(e)}, \quad i = 1, 2, \dots, \alpha_e$$

besteht eine unimodulare Transformation.

Beweis: Zunächst gilt

$$(2) \quad \begin{cases} \bar{a}_i^{(e)} = \sigma_{ij}^{(e)} a_j^{(e)} \\ a_i^{(e)} = \tau_{ij}^{(e)} \bar{a}_j^{(e)}, \quad i, j = 1, \dots, \alpha_e. \end{cases}$$

Daraus folgt $\bar{a}_i^{(q)} = \sigma_{ij}^{(q)} \tau_{jh}^{(q)} \bar{a}_h^{(q)}$, $i, j, h = 1, \dots, \alpha_q$, also $0^{(q)} = (\delta_{ih} - \sigma_{ij}^{(q)} \tau_{jh}^{(q)}) \bar{a}_h^{(q)}$, wo $\delta_{ih} = 0$ für $\begin{cases} i \neq h \\ i = h \end{cases}$ ist. Da aber auch $0^{(q)} = 0 \bar{a}_h^{(q)}$ gilt und jeder Komplex durch eine Basis eindeutig dargestellt wird, so folgt $\sigma_{ij}^{(q)} \tau_{jh}^{(q)} = \delta_{ih}$, q. e. d.

Der Randkomplex eines Komplexes $K^{(q)}$.

Wir führen ihn ein durch das

IV. Axiom: Es sei eine Operation „Randbildung“ vorhanden, die jedem ρ dimens. Komplex $K^{(q)}$ in Σ ($\rho = 1, 2, \dots, n$) einen bestimmten $\rho - 1$ dimens. Komplex von Σ , seinen Rand, $R(K^{(q)})$ bezeichnet, zuordnet.

Die Operation „Randbildung“ selbst genüge dabei den Axiomen:

V. Axiom: Seien t_1 und t_2 ganze Zahlen und $K_1^{(q)}, K_2^{(q)}$ zwei ρ dimens. Komplexe in Σ , so gelte:

$$(3) \quad \underline{R(t_1 K_1^{(q)} + t_2 K_2^{(q)}) = t_1 R(K_1^{(q)}) + t_2 R(K_2^{(q)})}$$

VI. Axiom: Der Randkomplex eines Randes sei ein Nullkomplex:

$$(4) \quad \underline{R(R(K^{(q)})) = 0^{(q-2)}}.$$

Infolge der Festsetzung (3) ist uns der Randkomplex $R(K^{(q)})$ eines beliebigen ρ dimens. Komplexes $K^{(q)}$ bekannt, sobald wir die Randkomplexe der Elemente einer Basis $a_i^{(q)}$, $i = 1, \dots, \alpha_q$ der ρ dimens. Komplexe kennen.

Sei

$$(5) \quad R(a_j^{(q)}) = \varepsilon_{ij}^{(q)} a_i^{(q-1)}, \quad j = 1 \dots \alpha_q, \quad i = 1, \dots, \alpha_{q-1},$$

$$\varepsilon_{ij}^{(q)} = \text{ganze Zahl},$$

das System dieser $\rho - 1$ dimens. Randkomplexe, dargestellt durch die Basis $a_i^{(q-1)}$, $i = 1 \dots \alpha_{q-1}$ von Σ . Aus

$$(6) \quad K^{(q)} = p_j a_j^{(q)}, \quad j = 1, \dots, \alpha_q$$

folgt dann nach (3) und (5)

$$(6') \quad R(K^{(q)}) = p_j \varepsilon_{ij}^{(q)} a_i^{(q-1)}, \quad j = 1 \dots \alpha_q, \quad i = 1, \dots, \alpha_{q-1}.$$

Wegen (4) folgt für $K_j^{(q)} = a_j^{(q)}$:

$$\begin{aligned} 0^{(q-2)} &= R(R(a_j^{(q)})) = R(\varepsilon_{ij}^{(q)} a_i^{(q-1)}) = \varepsilon_{ij}^{(q)} R(a_i^{(q-1)}) \\ &= \varepsilon_{ij}^{(q)} \varepsilon_{k i}^{(q-1)} a_k^{(q-2)}, \quad j = 1, \dots, \alpha_q, \quad i = 1, \dots, \alpha_{q-1}, \quad k = 1, \dots, \alpha_{q-2}, \end{aligned}$$

und da $a_k^{(q-2)}$, $k = 1 \dots \alpha_{q-2}$, Basis der $\rho-2$ dimens. Komplexe von Σ ist, so gilt demnach:

$$(7) \quad \sum_{i=1}^{\alpha_{q-1}} \varepsilon_{ij}^{(q)} \varepsilon_{k i}^{(q-1)} = 0, \quad j = 1, \dots, \alpha_q, \quad k = 1, \dots, \alpha_{q-2}$$

Ist (7) erfüllt, so ist $R(R(a^{(q)})) = 0^{(q-2)}$ und weiter $R(R(K^{(q)})) = 0^{(q-2)}$.

Bemerkung: Ist für ein Basissystem $a_i^{(q)}$, $\rho = 0, 1, \dots, n$ und $i = 1, \dots, \alpha_q$ das System der in (5) auftretenden Koeffizienten $\varepsilon_{ij}^{(q)}$ bekannt, so erhalten wir dieses Koeffizientensystem für ein neues Basissystem (2) $\bar{a}_i^{(q)}$, $\rho = 0, 1 \dots n$, $i = 1, \dots, \alpha_q$ aus:

$$\begin{aligned} R(\bar{a}_i^{(q)}) &= \bar{\varepsilon}_{hi}^{(q)} \bar{a}_h^{(q-1)} = \sigma_{ij}^{(q)} R(a_j^{(q)}) = \sigma_{ij}^{(q)} \varepsilon_{kj}^{(q)} a_k^{(q-1)} \\ &= \sigma_{ij}^{(q)} \varepsilon_{kj}^{(q)} \tau_{kh}^{(q-1)} \bar{a}_h^{(q-1)} \quad \text{in der Form:} \end{aligned}$$

$$(8) \quad \bar{\varepsilon}_{hi}^{(q)} = \sigma_{ij}^{(q)} \varepsilon_{kj}^{(q)} \tau_{kh}^{(q-1)}, \quad \begin{cases} h, k = 1, \dots, \alpha_{q-1} \\ i, j = 1, \dots, \alpha_q \end{cases}$$

Die ρ dimensionalen Zyklen:

Ein ρ dimens. Komplex, dessen Randkomplex der $\rho-1$ dimens. Nullkomplex ist, heie ein ρ dimens. Zyklus.

In dieser Definition ist $\rho > 0$ vorausgesetzt, nulldimensionale Komplexe seien stets als nulldimens. Zyklen bezeichnet.

Die Gesamtheit der ρ dimens. Zyklen in Σ bildet eine Untergruppe der Gruppe $\{K^{(q)}\}$ der ρ dimens. Komplexe in Σ , wir bezeichnen sie $\{C^{(q)}\}$.

Denn nach (3) ist die Summe zweier Zyklen ein Zyklus, und ist $C^{(q)}$ ein Zyklus, so auch $-C^{(q)}$. Soll $K^{(q)} = p_j a_j^{(q)}$, $j = 1, \dots, \alpha_q$, ein Zyklus sein, so mssen die ganzen Zahlen p_j , $j = 1, \dots, \alpha_q$, nach (6') das System linearer Gleichungen

$$(9) \quad p_j \varepsilon_{ij}^{(q)} = 0, \quad j = 1 \dots \alpha_q, \quad i = 1, \dots, \alpha_{q-1},$$

erfllen.

Mit r_q bezeichnen wir den Rang der Matrix $\|\varepsilon_{ij}^{(q)}\|$, dann hat

(9) $\alpha_q - r_q$ unabhngige Lsungssysteme $p_1, p_2, \dots, p_{\alpha_q}$.

Wir bestimmen nun eine Basis der ρ dimens. Zyklen, die nach dem Vorhergehenden $\alpha_q - r_q$ unabhngige Elemente hat und verwenden zu diesem Behufe das System (5). Durch unimodulare Trans-

formation der Basiselemente der ρ resp. $\rho-1$ dimens. Komplexe $\alpha_j^{(q)}$, $\alpha_i^{(q-1)}$ bringen wir jenes System in die kanonische Form:

$$(10) \quad R(b_j^{(q)}) = E_{ij}^{(q)} b_i^{(q-1)}, \quad j = 1, \dots, \alpha_q, \quad i = 1, \dots, \alpha_{q-1},$$

wo die $b_j^{(q)}$ eine neue Basis der ρ dimens. Komplexe in Σ ,
 die $b_i^{(q-1)}$ " " " " $\rho-1$ " " " "
 und die ganzen Zahlen:

$$(11) \quad E_{11}^{(q)}, E_{22}^{(q)}, \dots, E_{r_q r_q}^{(q)},$$

die Gesamtheit der von Null verschiedenen Koeffizienten des Systemes (10), die invarianten Faktoren der Matrix $\|E_{ij}^{(q)}\|$ sind. Nach (8) ist die Reihe (11) von der speziellen Basiswahl unabhängig, die von 1 verschiedenen Größen der Reihe (11) heißen die $\rho-1$ ten Torsionszahlen von Σ .

Für sie gilt, daß für $p < q$ $E_{pp}^{(q)}$ Faktor von $E_{qq}^{(q)}$ ist.

Aus (10) folgt, da $E_{ij}^{(q)} = 0$ ist für $i \neq j$ und $i > r_q$:

$$(12) \quad R(b_j^{(q)}) = 0^{(q-1)}, \quad j = r_q + 1, r_q + 2, \dots, \alpha_q,$$

das besagt aber, daß das Teilsystem der Basis $b_j^{(q)}$, $j = 1, 2, \dots, \alpha_q$

$$(13) \quad b_{r_q+1}^{(q)}, b_{r_q+2}^{(q)}, \dots, b_{\alpha_q}^{(q)}$$

$\alpha_q - r_q$ unabhängige ρ dimens. Zyklen sind.

Sei nun $C^{(q)}$ ein ρ dimens. Zyklus in Σ ; als ρ dimens. Komplex ist er jedenfalls darstellbar durch die Basis $b_j^{(q)}$, $j = 1, \dots, \alpha_q$:

$$(14) \quad C^{(q)} = p_j^q b_j^{(q)}, \quad j = 1, \dots, \alpha_q.$$

Die Operation „Randbildung“ ergibt

$$(14') \quad R(C^{(q)}) = 0^{(q-1)} = p_j E_{ij}^{(q)} b_i^{(q-1)}, \quad j = 1, \dots, \alpha_q, \quad i = 1, \dots, \alpha_{q-1}$$

und demnach $p_j E_{ij}^{(q)} = 0$, d. h. $p_1 E_{11}^{(q)} = 0$, $p_2 E_{22}^{(q)} = 0$, $\dots$,
 $p_{r_q} E_{r_q r_q}^{(q)} = 0$, also $p_1 = p_2 = \dots = p_r = 0$.

Demnach treten in der Darstellung (14) nur die Elemente $b_j^{(q)}$, $j = r_q + 1, \dots, \alpha_q$, auf, d. h. aber, daß die Reihe (13) eine Basis der ρ dimens. Zyklen von Σ darstellt.

Die Zyklen homolog Null.

Wegen $R(R(k^{(q+1)})) = 0^{(q-1)}$ ist der Randkomplex eines $\rho+1$ dimens. Komplexes ein ρ dimens. Zyklus.

Wir nennen einen ρ dimens. Zyklus $C^{(e)}$ in Σ homolog Null in Σ , und schreiben dies

$$(15) \quad C^{(e)} \sim 0 \text{ in } \Sigma,$$

wenn er Randkomplex eines $\rho + 1$ dimens. Komplexes $K^{(e+1)}$ von Σ ist.

$C^{(e)} \sim 0$ in Σ bedeutet also $C^{(e)} = R(k^{(e+1)})$, wo $k^{(e+1)}$ ein $\rho + 1$ dimens. Komplex in Σ ist. Die ρ dimens. Zyklen homolog Null bilden eine Untergruppe der ρ dimens. Zyklen $\{C^{(e)}\}$ von Σ . Zerlegt man die Gruppe $\{C^{(e)}\}$ der ρ dimens. Zyklen mittels dieser Untergruppe, so erhält man als Faktorgruppe die in der Topologie bedeutsame ρ^{te} Homologiegruppe. Wir wollen jedoch diese Gruppe ohne Bezugnahme auf andere Gruppen einführen. Zu diesem Behufe definieren wir:

Wir sagen zwei ρ dimens. Zyklen $C_1^{(e)}$ und $C_2^{(e)}$ sind homolog in Σ und schreiben

$$(16) \quad C_1^{(e)} \sim C_2^{(e)} \text{ in } \Sigma, \text{ wenn } C_1^{(e)} - C_2^{(e)} \sim 0 \text{ in } \Sigma \text{ ist.}$$

Aus $C_1^{(e)} \sim C_2^{(e)}$ und $C_2^{(e)} \sim C_3^{(e)}$ folgt dann $C_1^{(e)} \sim C_3^{(e)}$.

Demnach lassen sich die ρ dimens. Zyklen in Klassen zusammenfassen der Eigenschaft, daß die Zyklen einer Klasse untereinander homolog sind, aber nicht Zyklen verschiedener Klassen.

Seien $\mathfrak{K}^{(e)}: \{C_1^{(e)}, C_2^{(e)}, \dots\}$ und $\mathfrak{K}^{(e)}: \{\bar{C}_1^{(e)}, \bar{C}_2^{(e)}, \dots\}$ zwei dieser Klassen, es gelte also

$$C_i^{(e)} \sim C_k^{(e)} \text{ in } \Sigma, \text{ resp. } \bar{C}_i^{(e)} \sim \bar{C}_k^{(e)} \text{ in } \Sigma,$$

dann behaupten wir, daß die Gesamtheit der Zyklen $C_p^{(e)} + \bar{C}_q^{(e)}$ wieder eine Klasse bilden.

Denn aus $C_p^{(e)} - C_r^{(e)} = R(k_1^{(e+1)})$ und $\bar{C}_q^{(e)} - \bar{C}_s^{(e)} = R(k_2^{(e+1)})$ folgt $(C_p^{(e)} + \bar{C}_q^{(e)}) - (C_r^{(e)} + \bar{C}_s^{(e)}) = R(k_1^{(e+1)} + k_2^{(e+1)})$, q. e. d.

Somit lassen sich die Zyklenklassen addieren und die Addition ist kommutativ.

Bezeichnen wir mit $\mathfrak{K}_0^{(e)}$ die Klasse (Gruppe in diesem Fall) der ρ dimens. Zyklen homolog Null, so ist $\mathfrak{K}_0^{(e)}$ das Nullelement dieser Addition. Weiter gibt es zu jeder Klasse $\mathfrak{K}^{(e)}$ ihre inverse $-\mathfrak{K}^{(e)}$, für die $\mathfrak{K}^{(e)} + (-\mathfrak{K}^{(e)}) = \mathfrak{K}_0^{(e)}$ ist.

²⁾ Sei $\tilde{C}^{(e)}$ umgekehrt irgend ein Element der Klasse $C_p^{(e)} + \bar{C}_q^{(e)}$, so hat $\tilde{C}^{(e)}$ die Gestalt $C_p^{(e)} + \bar{C}_q^{(e)} + R(k^{(e+1)})$, ist also gleich $C_p^{(e)} + [\bar{C}_q^{(e)} + R(k^{(e+1)})] = C_p^{(e)} + \bar{C}_r^{(e)}$.

(Enthält $\mathfrak{R}^{(q)}$ den Zyklus $C^{(q)}$, so ist $-\mathfrak{R}^{(q)}$ die Gesamtheit der Zyklen homolog $-C^{(q)}$ in Σ .) Endlich gilt auch das assoziative Gesetz $(\mathfrak{R}_1^{(e)} + \mathfrak{R}_2^{(e)}) + \mathfrak{R}_3^{(e)} = \mathfrak{R}_1^{(e)} + (\mathfrak{R}_2^{(e)} + \mathfrak{R}_3^{(e)})$, denn ist $C_i^{(e)}$ ein Zyklus von $\mathfrak{R}_i^{(e)}$, $i = 1, 2, 3$, so ist $(C_1^{(e)} + C_2^{(e)}) + C_3^{(e)} = C_1^{(e)} + (C_2^{(e)} + C_3^{(e)})$ ein Zyklus der beiden Klassen $(\mathfrak{R}_1^{(e)} + \mathfrak{R}_2^{(e)}) + \mathfrak{R}_3^{(e)}$, $\mathfrak{R}_1^{(e)} + (\mathfrak{R}_2^{(e)} + \mathfrak{R}_3^{(e)})$, die demnach zusammenfallen.

Die so definierten Klassen der ρ dimens. Zyklen sind demnach Elemente einer kommutativen Gruppe mit der Addition als Verknüpfung und der Klasse (Gruppe) der ρ dimens. Zyklen homolog Null $-\mathfrak{R}_0-$ als Nullelement.

Die Gruppe hat den Namen ρ^{te} Homologiegruppe von Σ .

Wir zeigen nun, daß es eine Basis der ρ dimens. Zyklen homolog Null gibt.

Jeder $\rho + 1$ dimens. Komplex läßt sich mittels einer Basis $a_j^{(e+1)}$, $j = 1, \dots, \alpha_{e+1}$ darstellen, somit jeder ρ dimens. Zyklus homolog Null mittels der Randkomplexe

$$(17) \quad R(a_j^{(e+1)}) = \varepsilon_{ij}^{(e+1)} a_i^{(e)}, \quad j = 1, \dots, \alpha_{e+1}, \quad i = 1, \dots, \alpha_e.$$

Die Anzahl der linear unabhängigen dieser Randkomplexe ist demnach gleich der Anzahl der linear unabhängigen ρ dimens. Zyklen homolog Null und gleich r_{e+1} , dem Range der Matrix $\|\varepsilon_{ij}^{(e+1)}\|$.

Durch unimodulare Transformation der Basissysteme $a_j^{(e+1)}$, $j = 1, \dots, \alpha_{e+1}$ resp. $a_i^{(e)}$, $i = 1, \dots, \alpha_e$, bringen wir (17) in die kanonische Form:

$$(18) \quad R(d_j^{(e+1)}) = E_{ij}^{(e+1)} d_i^{(e)}, \quad j = 1, \dots, \alpha_{e+1}, \quad i = 1, \dots, \alpha_e,$$

wobei die $d_j^{(e+1)}$ resp. $d_i^{(e)}$ neue Basissysteme der $\rho + 1$ resp. ρ dimens. Komplexe in Σ und die ganzen Zahlen

$$(18') \quad E_{11}^{(e+1)}, E_{22}^{(e+1)}, \dots, E_{r_{e+1} r_{e+1}}^{(e+1)}, \text{ die Gesamtheit der von Null}$$

verschiedenen Koeffizienten des Systemes (18), die invarianten Faktoren der Matrix $\|\varepsilon_{ij}^{(e+1)}\|$ sind.

Wir behaupten, daß die r_{e+1} ρ dimens. Zyklen

$$(19) \quad R(d_j^{(e+1)}) = N_j^{(e)}, \quad j = 1, \dots, r_{e+1}$$

eine Basis der ρ dimens. Zyklen homolog Null bilden.

In der Tat läßt sich jeder $\rho + 1$ dimens. Komplex $K^{(e+1)}$ eindeutig mittels der Basis $d_j^{(e+1)}$, $j = 1, \dots, \alpha_{e+1}$ darstellen: $K^{(e+1)} = \sigma_j d_j^{(e+1)}$, $j = 1, \dots, \alpha_{e+1}$. Daher gilt $R(K^{(e+1)}) = \sigma_j R(d_j^{(e+1)}) =$

$= \sigma_j N_j^{(q)}$, $j=1, \dots, r_{q+1}$, da $R(d_j^{(q+1)})=0$ ist für $j > r_{q+1}$. Da aber $R(K^{(q+1)})$ der allgemeinste Zyklus homolog Null ist, so läßt sich jeder Zyklus ρ^{ter} Dimension homolog Null mittels der $N_j^{(q)}$, $j=1, \dots, r_{q+1}$ darstellen. Diese Darstellung ist weiter eindeutig, denn aus $\pi_j N_j^{(q)} = 0^{(q)}$, $j=1, \dots, r_{q+1}$ folgt nach (18) $\pi_1 E_{11}^{(q)} d_1^{(q)} + \pi_2 E_{22}^{(q)} d_2^{(q)} + \dots + \pi_{r_{q+1}} E_{r_{q+1} r_{q+1}}^{(q)} d_{r_{q+1}}^{(q)} = 0^{(q)}$, und da $d_j^{(q)}$, $j=1 \dots r_q$ eine Basis ist $\pi_1 E_{11}^{(q)} = \pi_2 E_{22}^{(q)} = \dots = \pi_{r_{q+1}} E_{r_{q+1} r_{q+1}}^{(q)} = 0$, somit $\pi_1 = \pi_2 = \dots = \pi_{r_{q+1}} = 0$.

Es gibt demnach kein System von weniger als r_{q+1} Elementen, durch das jeder ρ dimens. Zyklus homolog Null darstellbar wäre.

Wir notieren das Resultat: Wir können stets eine Basis der $\rho+1$ dimens. Komplexe in $\Sigma d_j^{(q+1)}$, $j=1, \dots, r_{q+1}$, so finden, daß $d_j^{(q+1)}$, $j=r_{q+1}+1, \dots, r_{q+1}+r_{q+1}$ eine Basis der $\rho+1$ dimens. Zyklen in Σ ist und $R(d_j^{(q+1)})$, $j=1, 2, \dots, r_{q+1}$ eine Basis der ρ dimens. Zyklen homolog Null in Σ ist.

Ferner vermerken wir: Die im kanonischen Systeme (18) auftretenden ρ dimens. Komplexe $d_i^{(q)}$, $i=1, 2, \dots, r_{q+1}$ sind, wie Randbildung zeigt, ρ dimens. Zyklen.

Für die den Zyklus $d_i^{(q)}$, $i=1, \dots, r_{q+1}$, enthaltende Klasse ist für $E_{ii}^{(q+1)} \neq 1$ charakteristisch, daß die Zyklen der Klasse nicht homolog Null sind, dagegen die mit $E_{ii}^{(q+1)}$ multiplizierten bereits diese Eigenschaft besitzen.

Daraus aber folgt, daß für Ringe Σ , für die die Größen der Reihe (18¹) nicht sämtlich eins sind, es keine Basis der Klassen der ρ dimens. Zyklen gibt³).

Denn wäre $\mathfrak{R}_1^{(q)}, \mathfrak{R}_2^{(q)}, \dots, \mathfrak{R}_\varepsilon^{(q)}$ eine solche Basis und $\mathfrak{R}^{(q)}$ eine Klasse der Eigenschaft $t\mathfrak{R}^{(q)} = \mathfrak{R}_0^{(q)}$, ($\mathfrak{R}_0^{(q)}$ das Nullelement), so folgt aus $\mathfrak{R}^{(q)} = k_i \mathfrak{R}_i^{(q)}$, $i=1, \dots, \varepsilon$ $t\mathfrak{R}^{(q)} = t k_i \mathfrak{R}_i^{(q)} = \mathfrak{R}_0^{(q)}$, woraus $t k_i = 0$, $i=1, \dots, \varepsilon$, also $k_i = 0$, d. h. $\mathfrak{R}^{(q)} = \mathfrak{R}_0^{(q)}$ geschlossen wird.

Dagegen gilt:

Die Elemente der ρ^{ten} Homologiegruppe werden durch eine Basis der ρ dimens. Zyklen (resp. der entsprechenden Klasse) dargestellt:

(20) $C_1^{(q)}, C_2^{(q)}, \dots, C_{r_q - r_q}^{(q)}$, wobei dann zwischen diesen Zyklen die definierenden Relationen bestehen:

$$(20') \quad N_j^{(q)} = \gamma_{ji} C_j^{(q)} \infty 0, \quad j=1, 2, \dots, r_{q+1}$$

³) Basis mit linear unabhängigen Elementen.

Die Differenz $\alpha_q - r_q - r_{q+1}$ der Anzahl der Basiselemente der ρ dimens. Zyklen und der der ρ dimens. Zyklen homolog Null

$$(21) \quad h_q = \alpha_q - r_q - r_{q+1}$$

heißt die ρ^{te} Bettische Zahl⁴⁾.

Sei nun

$$(22) \quad C_1^{(q)}, C_2^{(q)}, \dots, C_{\alpha_q - r_q}^{(q)} \quad \text{eine Basis der } \rho \text{ dimens. Zyklen und}$$

$$(23) \quad M_i^{(q)} = \tau_{ji} C_j^{(q)} \infty 0, \quad i = 1, \dots, r_{q+1}, \quad j = 1, \dots, \alpha_q - r_q$$

eine Basis der ρ dimens. Zyklen homolog Null, so bringen wir (23) durch unimodulare Transformationen der beiden Basissysteme in die kanonische Form:

$$(24) \quad \overline{M}_i^{(q)} = \overline{E}_{ji} \overline{C}_j^{(q)}, \quad i = 1, \dots, r_{q+1}, \quad j = 1, \dots, \alpha_q - r_q, \quad \text{wo also}$$

die von Null verschiedenen Koeffizienten $\overline{E}_{ji}$:

$$(24') \quad \overline{E}_{11}, \overline{E}_{22}, \dots, \overline{E}_{r_{q+1} r_{q+1}}, \quad \text{die invarianten Faktoren der}$$

Matrix $\|\tau_{ji}\|$ sind. (Der Rang dieser Matrix kann nicht kleiner als r_{q+1} sein, da es sonst nach (23) eine Beziehung $\pi_i M_i^{(q)} = 0^{(q)}$, $i = 1, \dots, r_{q+1}$, gäbe mit von Null verschiedenen π_i , im Widerspruch mit der Basiseigenschaft von $M_i^{(q)}$, $i = 1, \dots, r_{q+1}$.)

Wir behaupten nun die Identität der beiden Reihen (18') und (24').

Da die $N_j^{(q)}$ und die $\overline{M}_j^{(q)}$ (24), $j = 1, \dots, r_{q+1}$ Basissysteme der ρ dimens. Zyklen homolog Null sind, so gilt:

$$(25) \quad \begin{cases} \overline{M}_j^{(q)} = \sigma_{jk} N_k^{(q)} \\ N_j^{(q)} = \tau_{jk} \overline{M}_k^{(q)}, \quad k, j = 1, \dots, r_{q+1}, \end{cases}$$

wo $\|\sigma_{jk}\|$ eine unimodulare Matrix ist und $\|\tau_{jk}\|$ ihre inverse. Dies ergibt nach (18), (19) und (24):

$$(25') \quad \begin{cases} \overline{E}_{ij} \overline{C}_i^{(q)} = \sigma_{jk} E_{ek}^{(q+1)} d_e^{(q)} \\ E_{ij}^{(q+1)} d_i^{(q)} = \tau_{jk} \overline{E}_{ek} \overline{C}_e^{(q)}, \quad \text{alle Indizes von 1 bis } r_{q+1}. \end{cases}$$

⁴⁾ Des öfteren wird $\alpha_q - r_q - r_{q+1} + 1$ als ρ^{te} Bettische Zahl eingeführt.

Nun sind die $d_i^{(e)}$ für $i = 1, \dots, r_{e+1}$, ρ dimens. Zyklen, somit darstellbar mittels der Basis $\bar{C}_1^{(e)}, \dots, \bar{C}_{\alpha_e - r_e}^{(e)}$, die $\bar{C}_i^{(e)}$, $i = 1, \dots, r_{e+1}$, wieder als ρ dimens. Komplexe lassen sich durch die Basis $d_1^{(e)}, d_2^{(e)}, \dots, d_{\alpha_e}^{(e)}$ darstellen. Nach (25') sind dann die $\bar{C}_i^{(e)}$, $i = 1, \dots, r_{e+1}$ allein durch die $d_i^{(e)}$, $i = 1, \dots, r_{e+1}$ und umgekehrt die $d_i^{(e)}$, $i = 1, \dots, r_{e+1}$ durch die $\bar{C}_i^{(e)}$, $i = 1, \dots, r_{e+1}$ allein darstellbar. Es besteht somit zwischen den $\bar{C}_i^{(e)}$ und den $d_i^{(e)}$, $i = 1, \dots, r_{e+1}$ eine unimodulare Transformation.

$$(26) \quad \bar{C}_e^{(e)} = \varepsilon_{ei} d_i^{(e)}, \quad i, e = 1, \dots, r_{e+1}.$$

Dies in die zweite Relation (25') eingesetzt, hat

$$(27) \quad E_{ij}^{(e+1)} d_i^{(e)} = \tau_{jk} \bar{E}_{ek} \varepsilon_{ei} d_i^{(e)}, \text{ also:}$$

$$(27') \quad E_{ij}^{(e+1)} = \tau_{jk} \bar{E}_{ek} \varepsilon_{ei}, \text{ alle Indizes von 1 bis } r_{e+1}, \text{ zur Folge.}$$

Da die $\|\tau_{jk}\|$ und die $\|\varepsilon_{ze}\|$ unimodulare Matrizen sind, so folgt hieraus die Identität der invarianten Faktoren, d. h. die Identität der beiden Reihen (18') und (24').

Hat eine Basis der ρ dimens. Zyklen in Σ

$$(28) \quad C_1^{(e)}, C_2^{(e)}, \dots, C_{\alpha_e - k_e}^{(e)}$$

zum System der definierenden Relationen

$$(29) \quad N_1^{(e)} = E_{11}^{(e+1)} C_1^{(e)} \infty 0, N_2^{(e)} = E_{22}^{(e+1)} C_2^{(e)} \infty 0, \dots$$

$$N_{r_{e+1}}^{(e)} = E_{r_{e+1} r_{e+1}}^{(e+1)} C_{r_{e+1}}^{(e)} \infty 0,$$

so heiße sie kanonisch.

Ist $C_i^{(e)}$ Element einer kanonischen Basis der ρ dimens. Zyklen, so folgt aus $E C_i^{(e)} \infty 0$ $E = p E_{ii}^{(e+1)}$. Denn $E C_i^{(e)}$ läßt sich dann durch die Basis (29), $N_1^{(e)}, \dots, N_{r_{e+1}}^{(e)}$ der ρ dimens. Zyklen homolog Null darstellen, wegen der Unabhängigkeit der $C_j^{(e)}$, $j = 1, \dots, \alpha_e - r_e$ kann diese Darstellung nur $E C_i^{(e)} = p E_{ii}^{(e+1)} C_i^{(e)}$ lauten, woraus $E = p E_{ii}^{(e+1)}$ folgt.

Man kann nach (29) die ρ^{ten} Torsionszahlen [Poincarésche Zahlen der ρ^{ten} Homologiegruppe⁵⁾] als das Koeffi-

⁵⁾ Poincarésche Zahlen nach Tietze, Monatshefte f. M. Ph. 19.

zientensystem der definierenden Relationen einer kanonischen Basis bezeichnen, wobei aber nur die von 1 verschiedenen Koeffizienten berücksichtigt werden.

Sei $C^{(e)}$ ein ρ dimens. Zyklus, so wollen wir die Zyklenklasse, die ihn enthält, mit $[C^{(e)}]$ bezeichnen.

Sei dann (28) wieder eine kanonische Basis; zwischen den $h_e = \alpha_e - r_e - r_{e+1}$ Klassen

$$(30) \quad [C_{r_{e+1}+1}^{(e)}], [C_{r_{e+1}+2}^{(e)}], \dots [C_{\alpha_e - r_e}^{(e)}]$$

besteht keine lineare Beziehung

$$(30') \quad A_{r_{e+1}+1} [C_{r_{e+1}+1}^{(e)}] + \dots + A_{\alpha_e - r_e} [C_{\alpha_e - r_e}^{(e)}] = \mathfrak{R}_e^{(e)},$$

denn diese hätte

$$\sum A_i C_i^{(e)} \propto 0, \quad i = r_{e+1} + 1, \dots, \alpha_e - r_e, \text{ zur Folge; das}$$

hieß aber

$$\sum_{r_{e+1}+1}^{\alpha_e - r_e} A_i C_i^{(e)} = \sum_1^{r_{e+1}} B_k N_k^{(e)}, \text{ woraus } A_i = 0, i = r_{e+1} + 1, \dots, \alpha_e - r_e$$

zu schließen ist.

Ist dagegen $[C^{(e)}]$ irgend eine weitere ρ dimens. Zyklenklasse, so folgt aus

$$\begin{aligned} [C^{(e)}] &= a_j [C_j^{(e)}], j = 1, \dots, \alpha_e - r_e, E_{r_{e+1} r_{e+1}}^{(e)} [C^{(e)}] = \\ &= \sum_{r_{e+1}+1}^{\alpha_e - r_e} E_{r_{e+1} r_{e+1}}^{(e)} a_j [C_j^{(e)}] = \mathfrak{R}_0^{(e)}. \end{aligned}$$

Wir können daher die ρ^{te} Bettische Zahl als größte Anzahl unabhängiger Zyklenklassen ρ dimens. Zyklen definieren.

Der Komplexring $\Sigma \pmod{2}^e$.

Sei (1') wieder eine Basis der ρ dimens. Komplexe in Σ .

Die Komplexe $K^{(e)} = \sigma_j a_j^{(e)}$, $\sigma_j \equiv 0 \pmod{2}$ bilden eine Untergruppe der Gruppe $\{K^{(e)}\}$. Zerlegen wir die Gruppe $\{K^{(e)}\}$ vermittle dieser Untergruppe, so erhalten wir eine Faktorgruppe, deren

^{e)} Komplexe mod 2 wurden zuerst von O. Veblen betrachtet.

Elemente jene Komplexe $K^{(e)} = \tau_j a_j^{(e)}$ vereinigt, deren Darstellungszahlen mod 2 kongruent sind.

Wir gelangen zu dieser Gruppe auch in folgender Weise:

Wir identifizieren zwei Komplexe

$$(31) \quad K^{(e)} = \sigma_j a_j^{(e)}, \quad K^{(e)} = \bar{\sigma}_j a_j^{(e)}, \quad j = 1, \dots, \alpha_e \text{ und schreiben}$$

$$(32) \quad K^{(e)} \equiv \bar{K}^{(e)}, \text{ wenn } \sigma_j \equiv \bar{\sigma}_j \pmod{2} \text{ ist.}$$

Aus $K_1^{(e)} \equiv K_2^{(e)}$ und $K_2^{(e)} \equiv K_3^{(e)}$ folgt dann $K_1^{(e)} \equiv K_3^{(e)}$.

Diese Identifizierung ist von der Wahl des Basissystems unabhängig.

Ist $\bar{a}_1^{(e)}, \dots, \bar{a}_{\alpha_e}^{(e)}$ eine zweite Basis, so besteht zwischen dieser und (1') der Zusammenhang

$$(33) \quad \begin{cases} \bar{a}_i^{(e)} \equiv \sigma_{ij}^{(e)} a_j^{(e)} \pmod{2} \\ a_i^{(e)} \equiv \tau_{ij}^{(e)} \bar{a}_j^{(e)} \pmod{2}, \end{cases}$$

wobei $\|\sigma_{ij}^{(e)}\|$ eine unimodulare Matrix ist und $\|\tau_{ij}^{(e)}\|$ ihre inverse.

Aus

$$(34) \quad \varepsilon_j a_j^{(e)} \equiv 0^{(e)} \text{ folgt } \varepsilon_j \equiv 0 \pmod{2}, \quad j = 1, \dots, \alpha_e, \text{ wenn } a_j^{(e)} \text{ eine}$$

Basis der ρ dimens. Zyklen ist.

Ferner gelten die Relationen:

$$\text{Aus } K^{(e)} \equiv K_1^{(e)} \text{ und } L^{(e)} \equiv L_1^{(e)} \text{ folgt}$$

$$t K^{(e)} + l L^{(e)} \equiv t K_1^{(e)} + l L_1^{(e)}.$$

Weiter hat $K^{(e)} \equiv K_1^{(e)}$ die Relation $R(K^{(e)}) \equiv R(K_1^{(e)})$ zur Folge.

Der Komplexring $\bar{\Sigma} \pmod{2}$ wird gebildet durch die Gesamtheit aller Komplexe $K^{(e)}$ in Σ , wobei die Gesamtheit identischer Komplexe als ein Element gezählt wird.

Ein Zyklus von $\bar{\Sigma}$ (unorientierter Zyklus) ist ein Komplex, dessen Randkomplex dem entsprechend dimensional Nullkomplex identisch ist.

Jeder Zyklus von Σ ist dann Zyklus von $\bar{\Sigma}$, aber nicht umgekehrt.

Betrachten wir die Basis (10) $b_1^{(e)}, b_2^{(e)}, \dots, b_{\alpha_e}^{(e)}$ der ρ dimens. Komplexe in Σ und ist $E_{s_e+1 s_e+1}^{(e)} \equiv 0, \pmod{2}$, dagegen nicht $E_{s_e s_e}^{(e)}$, so ist

$$(35) \quad b_{s_0+1}^{(e)}, b_{s_0+2}^{(e)}, \dots, b_{\alpha_0}^{(e)}$$

eine Basis der ρ dimens. Zyklen von $\bar{\Sigma}$.

Ungeändert definieren wir: Die Randkomplexe $\rho + 1$ dimens. Komplexe heißen ρ dimens. Zyklen homolog Null.

Wir erhalten eine Basisdarstellung der ρ dimens. Zyklen homolog Null von Σ aus (18), wobei genau jene $N_j^{(e)}$ auftreten, für die $E_{jj}^{e+1} \not\equiv 0, \text{ mod } 2$, ist.

Dies gelte bis $j = s_0 + 1$, dann ist

$$(36) \quad N_1^{(e)}, N_2^{(e)}, \dots, N_{s_0+1}^{(e)} \quad \text{eine solche Basis.}$$

Die der Bettischen Zahl in $\bar{\Sigma}$ entsprechende $\alpha_0 - s_0 - s_0 + 1$ heißt ρ^{te} Zusammenhangszahl, die der ρ^{ten} Homologiegruppe entsprechende Gruppe ρ^{te} Zusammenhangsgruppe.

II. ABSCHNITT.

Unterring (Teilring) eines Ringes Σ .

Ein Teilsystem der Komplexe des Ringes Σ , das den Axiomen I bis VI genügt, heie Unterring von Σ , und zwar echter Unterring, wenn es in Σ Komplexe gibt, die in diesem Teilsystem nicht enthalten sind.

Wir betrachten in der Folge Unterringe, die in bezug auf ein fest gewhltes Basissystem

$$(37) \quad a_i^{(e)}, \quad i = 1, \dots, \alpha_0, \quad \rho = 0, 1, \dots, n$$

des Ringes Σ definiert sind. Entnimmt man einem solchen Basissystem ein Teilsystem S , so bildet die Gesamtheit der durch die Elemente von S dargestellten Komplexe von Σ im allgemeinen keinen Ring, da ja der Rand eines solchen Komplexes in dieser Gesamtheit nicht enthalten sein mu.

Durch Adjunktion aller Basiselemente, die Rnder sind der Elemente von S , aber nicht zu S gehren, erweitert man das System S und gelangt nach endlich vielen solchen Erweiterungen zum kleinsten Ringe ber S .

Ein n dimens. Ring Σ , der der kleinste Ring ist ber dem System seiner n dimensional Basiselemente, heit homogen (homogen in bezug auf die Dimension).

Sei m eine positive ganze Zahl kleiner oder gleich n , wo n die Dimension eines Ringes Σ ist und sei (37) ein Basissystem von Σ .

Die durch die Basiselemente

$$(38) \quad a_i^{(e)}, \quad i = 1, \dots, \alpha_0, \quad \rho = 0, 1, \dots, m,$$

dargestellten Komplexe bilden einen m dimens. Unterring $\Gamma^{(m)}$ von Σ , den man das m dimens. Gerüst von Σ nennt. Es ist $\Gamma^{(m)}$ mit Σ identisch.

Ist Σ ein homogener Ring, so ist auch $\Gamma^{(m)}$ homogen.

Das m dimens. Gerüst wird durch die Gesamtheit der Komplexe von Σ gebildet, deren Dimension kleiner ist als $m + 1$, somit ist $\Gamma^{(m)}$ von der Wahl des Basissystems (37) von Σ unabhängig.

Der Komplementärring und der Randring eines echten Unterringes Σ_1 von Σ .

Es sei Σ_1 ein durch ein Teilsystem von (37) definierter Unterring der Dimension $m \leq n$ von Σ . (Die Definition des Komplementärringes resp. Randringes wird nur für solche Unterringe gegeben.) Σ_1 ist dann auch Unterring von $\Gamma^{(m)}$. Streichen wir in (38) alle Basiselemente von Σ_1 , so wird durch das Restsystem S eine Gesamtheit von Komplexen dargestellt, die $\Gamma^{(m)} - \Sigma_1$ bezeichnet sei und die im allgemeinen keinen Ring bildet.

Als Komplementring Σ_2 von Σ_1 definieren wir dann den kleinsten Ring über S .

Ist Σ ein homogener Ring und Σ_1 ebenfalls homogen, so ist auch Σ_2 homogen. Beweis: Von den m dimens. Basiselementen von $\Gamma^{(m)}$ $a_i^{(m)}$, $i = 1, \dots, \alpha_m$, mögen die Elemente $a_i^{(m)}$, $i = 1, \dots, \beta$, ($\beta \leq \alpha_m$) die m dimens. Basiselemente von Σ_1 sein, $a_i^{(m)}$, $i = \beta + 1, \dots, \alpha_m$ sind dann die m dimens. Basiselemente von Σ_2 . Wäre nun Σ_2 nicht mit dem kleinsten Ringe über $a_i^{(m)}$, $i = \beta + 1, \dots, \alpha_m$, identisch, so müßte es in Σ_2 ein Element $a_j^{(h)}$, $h < m$, geben, das in diesem kleinsten Ringe nicht enthalten wäre, also — da Σ homogen — im kleinsten Ringe über $a_i^{(m)}$, $i = 1 \dots \beta$, d. h. in Σ_1 enthalten sein müßte.

Gemeinsam aber sind Σ_1 und Σ_2 nur solche Elemente $a_j^{(h)}$, die im Rand eines Elementes von Σ_2 der Dimension $h + 1$ enthalten sind. Sei $a_k^{(h+1)}$ dieses Element von Σ_2 , so könnte es ebenfalls nicht im kleinsten Ring über $a_i^{(m)}$, $i = \beta + 1, \dots, \alpha_m$, enthalten sein, denn sonst wäre auch $a_j^{(h)}$ in diesem kleinsten Ringe enthalten. Unter den Elementen von Σ_2 , die in diesem kleinsten Ring nicht enthalten sind, gibt es also keine Elemente höchster Dimension, d. h. es gibt überhaupt keine solchen Elemente.

Ist also Σ_1 ein m dimens. homogener Unterring im homogenen Ringe Σ der Dimension n , so ist der Komplementring Σ_2 , falls er nicht leer ist ($\Gamma^{(m)} \equiv \Sigma_1$), ebenfalls homogen von der Dimension m . Die Beziehung zwischen Ring und Komplementring ist in diesem Falle wechselseitig.

Ohne weiteres ist klar, daß der Durchschnitt zweier Unterringe von Σ selbst ein Unterring von Σ ist.

Insbesondere nennt man den Durchschnitt des Unterringes Σ_1 und seines Komplementes Σ_2 den Randring Σ_r von Σ_1 .

Ist Σ_1 ein homogener Unterring im homogenen Ringe Σ , so ist Σ_r , da Σ_1 dann auch das Komplement von Σ_2 ist, gleichfalls Randring von Σ_2 .

Wir wollen in der Folge fast nur homogene Ringe betrachten⁷⁾, für sie gilt also: Der Durchschnitt (Σ_1, Σ_2) zweier Komplementringe Σ_1 und Σ_2 in Σ heißt Randring dieser Ringe. Sind Σ_1 und Σ_2 von der Dimension m , so ist Σ_r höchstens von der Dimension $m-1$.

Wir beweisen zwei für die Folge zu benutzende Hilfsätze:

Hilfssatz I: Ist Σ ein n dimens. Komplexring und Σ_1 ein Unterring von Σ , dessen Basissystem ein Teilsystem des Basissystems (37) von Σ ist, dann gibt es stets eine Basisdarstellung der ρ dimens. Zyklen von Σ_1 , die Teilsystem ist einer Basis der ρ dimens. Zyklen von Σ .

Beweis: Sei $C_1^{(e)}, C_2^{(e)}, \dots, C_s^{(e)}$ eine Basis der ρ dimens. Zyklen von Σ und $D_1^{(e)}, D_2^{(e)}, \dots, D_t^{(e)}$ eine Basis der ρ dimens. Zyklen von Σ_1 .

Da jeder Zyklus von Σ_1 auch Zyklus von Σ ist, gilt

$$D_i^{(e)} = a_{ji} C_j^{(e)}, i = 1, \dots, t, j = 1, \dots, s.$$

Dieses System bringen wir durch unimodulare Transformationen in die kanonische Form:

$$\overline{D}_1^{(e)} = a_1 \overline{C}_1^{(e)}, \dots, \overline{D}_t^{(e)} = a_t \overline{C}_t^{(e)}, \text{ wo } \overline{D}_i^{(e)}, i = 1, \dots, t \text{ eine neue Basis}$$

der ρ dimens. Zyklen von Σ_1 und $\overline{C}_j^{(e)}, j = 1, \dots, s$, die entsprechende Basis für Σ darstellt.

Die Zyklen $\overline{C}_1^{(e)}, \dots, \overline{C}_t^{(e)}$ liegen schon im Ringe Σ_1 , denn ist $a_1^{(e)}, \dots, a_{\alpha_e}^{(e)}$ die Basis der ρ dimens. Komplexe von Σ ; $a_1^{(e)}, \dots, a_{\beta}^{(e)}, \beta < \alpha_e$, die entsprechende Basis von Σ_1 , so gilt $\overline{C}_i^{(e)} = \gamma_{ik} a_k^{(e)}, k = 1, \dots, \alpha_e, i = 1, \dots, t$.

Ebenso gilt

$$\overline{D}_i^{(e)} = a_i \overline{C}_i^{(e)} = \varepsilon_{ie} a_e^{(e)}, e = 1, \dots, \beta, i = 1, \dots, t$$

(über i keine Summation).

Aus den beiden letzten Gleichungen aber folgt $\gamma_{ik} = 0, i = 1, \dots, t, k = \beta + 1, \dots, \alpha_e$, d. h. $\overline{C}_i^{(e)}, i = 1, \dots, t$, liegt bereits in Σ_1 .

⁷⁾ Wir schreiben fast nur, denn der Durchschnitt homogener Unterringe des homogenen Ringes Σ , beispielsweise der Randring eines homogenen Ringes, muß selbst kein homogener Ring sein.

Da $D_1^{(e)}, \dots, \bar{D}_t^{(e)}$ eine Basis der ρ dimens. Zyklen von Σ_1 ist, gilt

$$C_h^{(e)} = b_{ih} D_i^{(e)}, \quad i, h = 1, \dots, t.$$

Multipliziert man diese Gleichung mit a_h , ohne über h zu summieren, so erhält man

$$\bar{D}_h^{(e)} = a_h b_{ih} \bar{D}_i^{(e)}, \quad i, h = 1, \dots, t,$$

woraus infolge der Basiseigenschaft von $D_i^{(e)}$, $b_{ih} = 0$ für $i \neq h$ und $a_h b_{hh} = 1$, also $a_h = b_{hh} = \pm 1$ folgt.

Also gilt bei entsprechender Vorzeichenwahl: $\bar{D}_1^{(e)} = \bar{C}_1^{(e)}$, $\bar{D}_2^{(e)} = \bar{C}_2^{(e)}, \dots, \bar{D}_t^{(e)} = \bar{C}_t^{(e)}$, q. e. d.

Der Hilfssatz I gilt nicht, sobald man statt Basis der ρ dimens. Zyklen schreibt: Basis der ρ dimens. Zyklen homolog Null. Denn sei

$N_1^{(e)}, N_2^{(e)}, \dots, N_\sigma^{(e)}$ eine Basis der ρ dimens. Zyklen homolog Null von Σ und

$M_1^{(e)}, M_2^{(e)}, \dots, M_\tau^{(e)}$ eine Basis der ρ dimens. Zyklen homolog Null von Σ_1 , so gilt, da jeder Zyklus von Σ_1 homolog Null in Σ_1 auch Zyklus von Σ homolog Null in Σ ist:

$M_i^{(e)} = c_{ij} N_j^{(e)}$, $i = 1, \dots, \tau$, $j = 1, \dots, \sigma$, welches System wir in kanonische Form bringen

$\bar{M}_1^{(e)} = c_1 N_1^{(e)}, \dots, \bar{M}_\tau^{(e)} = c_\tau N_\tau^{(e)}$. Hier aber können wir nur behaupten, daß $\bar{N}_j^{(e)}$, $j = 1, \dots, \tau$, ein ρ dimens. Zyklus von Σ_1 ist, aber nicht, daß er homolog Null in Σ_1 ist.

Wir nennen nun einen m dimens. Ring Σ_1 **Nullring**, wenn jeder ρ dimens. Zyklus von Σ_1 , für $\rho = 1, \dots, m-1$, ein Zyklus homolog Null in Σ_1 ist.

Dann aber gilt für die im Hilfssatze I besprochenen Ringe der

Hilfssatz II: Ist Σ ein n dimens. Komplexring und Σ_1 sein m dimens. Unterring, ein Nullring, so gibt es für $0 < \rho < m$ stets eine Basis der ρ dimens. Zyklen von Σ homolog Null in Σ , die als Teilsystem eine Basis der ρ dimens. Zyklen von Σ_1 homolog Null in Σ_1 enthält.

Denn jetzt ist $\bar{N}_j^{(e)}$, $j = 1, \dots, \tau$, da $0 < \rho < m$, als ρ dimens. Zyklus von Σ_1 auch homolog Null in Σ_1 und demnach linear durch die $\bar{M}_j^{(e)}$, $j = 1, \dots, \tau$ darstellbar. Der weitere Schluß verläuft dann wie beim Hilfssatze I.

Der reziproke, duale Ring eines Ringes.

Ist Σ ein gegebener n dimens. Ring und (37) ein Basissystem von Σ , so kann man Σ auf mancherlei Arten Ringe zuordnen.

Setzt man z. B. ($h < n$): $a_i^{(0)} = b_i^{(0-h)}$, $i = 1, \dots, \alpha_0$, $\rho = h, \dots, n$ und definiert weiter $R(a_i^{(0)}) = R(b_i^{(0-h)}) = \varepsilon_{ij}^{(0)} b_j^{(0-h-1)}$, so kann man die $b_i^{(\sigma)}$, $\sigma = 0, \dots, n-h$ als Basissystem eines $n-h$ dimens. Ringes auffassen.

Unter allen einem gegebenen Ringe zugeordneten Ringen hat für die Anwendung der reziproke (oder duale) Ring eine gewisse Bedeutung. Wir bringen im folgenden seine Definition und leiten hierauf einen Satz ab für geometrische Mannigfaltigkeiten, die man durch Komplexringe „approximieren“ kann.

Sei also Σ der gegebene n dimens. Ring und

$$(37) \quad a_i^{(0)}, \quad i = 1, \dots, \alpha_0, \quad \rho = 0, 1, \dots, n$$

ein Basissystem von Σ . Für dieses Basissystem gelten die Randrelationen

$$(39) \quad R(a_i^{(h)}) = \varepsilon_{ji}^{(h)} a_j^{(h-1)}, \quad i = 1, \dots, \alpha_h, \quad j = 1, \dots, \alpha_{h-1}, \text{ wobei}$$

$$(39') \quad \varepsilon_{ij}^{(0)} \varepsilon_{ki}^{(0-1)} = 0 \text{ ist, } j = 1, \dots, \alpha_0, \quad i = 1, \dots, \alpha_{0-1}, \quad k = 1, \dots, \alpha_{0-2}.$$

Wir definieren nun den reziproken Ring Σ^* von Σ durch die Basis

$$(40) \quad A_i^{(n-h)} = a_i^{(h)}, \quad i = 1, \dots, \alpha_h, \quad h = 0, \dots, n,$$

die den Randrelationen

$$(41) \quad R(A_i^{(n-h)}) = \varepsilon_{ij}^{(h+1)} A_j^{(n-h-1)}, \quad i = 1, \dots, \alpha_h, \quad j = 1, \dots, \alpha_{h+1}$$

genügen.

In der Tat ist dann:

$$(42) \quad \begin{aligned} R(R(A_i^{(n-h)})) &= \varepsilon_{ij}^{(h+1)} R(A_j^{(n-h-1)}) \\ &= \varepsilon_{ij}^{(h+1)} \varepsilon_{jk}^{(h+2)} A_k^{(n-h-2)} = 0^{(n-h-2)}, \quad i = 1, \dots, \alpha_h, \\ &\quad j = 1, \dots, \alpha_{h+1}, \quad k = 1, \dots, \alpha_{h+2} \end{aligned}$$

da ja $\varepsilon_{jk}^{(h+2)} \varepsilon_{ij}^{(h+1)} = 0$ ist. ($k = 1, \dots, \alpha_{h+2}$, $j = 1, \dots, \alpha_{h+1}$, $i = 1, \dots, \alpha_h$)

Die Beziehung Σ^* zu Σ ist wechselseitig.

Nach (41) ist, wenn wir wie im Abschnitt I $R(A_i^{(n-h)}) = \varepsilon_{ji}^{*(n-h)}$ $A_j^{(n-h-1)}$, also $\varepsilon_{ij}^{(h+1)} = \varepsilon_{ji}^{*(n-h)}$ bezeichnen, ferner mit α_{n-h}^* die Basiszahl der $n-h$ dimens. Komplexe und mit r_{n-h}^* den Rang der Matrix $\|\varepsilon_{ji}^{*(n-h)}\|$, die $n-h$ te Bettische Zahl von Σ^*

$$(43) \quad h_{n-h}^* = \alpha_{n-h}^* - r_{n-h}^* - r_{n-h+1}^*.$$

Nun ist $\alpha_{n-h}^* = \alpha_h$, $r_{n-h}^* = r_{h+1}$, $r_{n-h+1}^* = r_h$, somit gilt

$$(43') \quad h_{n-h}^* = \alpha_h - r_h - r_{h+1} = h_h.$$

Die h te Bettische Zahl von Σ ist gleich der $n-h$ ten Bettischen Zahl von Σ^* und umgekehrt.

Ferner sind die Torsionszahlen der ρ ten Homologiegruppe von Σ^* die von 1 verschiedenen invarianten Faktoren der Matrix $\|\varepsilon_{ij}^{*(\rho+1)}\|$, also der Matrix $\|\varepsilon_{ij}^{(n-\rho)}\|$ und somit identisch mit den Torsionszahlen der $n-\rho-1$ ten Homologiegruppe von Σ .

Die Torsionszahlen der ρ ten Homologiegruppe von Σ^* sind mit den Torsionszahlen der $n-\rho-1$ ten Homologiegruppe von Σ identisch.

Der von uns definierte reziproke Ring Σ^* ist von der Wahl der Basis (37) von Σ nicht unabhängig. Sei nämlich

$$(44) \quad \bar{a}_i^{(h)}, i=1, \dots, \alpha_h, h=0, \dots, n \text{ ein zweites Basissystem und}$$

$$(45) \quad \bar{a}_i^{(h)} = \sigma_{ij}^{(h)} a_j^{(h)}, \quad a_i^{(h)} = \tau_{ij}^{(h)} \bar{a}_j^{(h)}, \quad i, j=1, \dots, \alpha_h$$

die unimodulare Beziehung zwischen beiden Basissystemen.

Der in bezug auf die Basis $\bar{a}_i^{(h)}$ definierte reziproke Ring $\bar{\Sigma}^*$ hat das Basissystem

$$(46) \quad \bar{A}_i^{(n-h)} = \bar{a}_i^{(h)}, \quad i=1, \dots, \alpha_h, h=0, \dots, n$$

und zwischen $A_i^{(n-h)}$ (40) und $\bar{A}_i^{(n-h)}$ bestehen zufolge (45) die unimodularen Transformationen

$$(47) \quad \begin{cases} \bar{A}_i^{(n-h)} = \sigma_{ij}^{(h)} A_j^{(n-h)} = \sigma_{ij}^{*(n-h)} A_j^{(n-h)} \\ A_i^{(n-h)} = \tau_{ij}^{(h)} \bar{A}_j^{(n-h)} = \tau_{ij}^{*(n-h)} \bar{A}_j^{(n-h)}, \end{cases}$$

wo $\sigma_{ij}^{(h)} = \sigma_{ij}^{*(n-h)}$ und $\tau_{ij}^{(h)} = \tau_{ij}^{*(n-h)}$ gesetzt wurde.

Für den Ring $\bar{\Sigma}^*$ gelten analog (41) die Randrelationen:

$$(48) \quad R(\bar{A}_i^{(n-h)}) = \bar{\varepsilon}_{ij}^{(h+1)} \bar{A}_j^{(n-h-1)},$$

wo zwischen den $\bar{\varepsilon}_{ij}^{(h+1)}$ und den $\varepsilon_{ij}^{(h+1)}$ — in (48) — die Relationen (8) bestehen:

$$(49) \quad \bar{\varepsilon}_{pq}^{(h+1)} = \sigma_{qr}^{(h+1)} \varepsilon_{sr}^{(h+1)} \tau_{sp}^{(h)}.$$

Faßt man dagegen nach (47) $A_i^{(n-h)}$, $i=1, \dots, \alpha_h$, $h=0, \dots, n$ auf als Basissystem des Ringes Σ^* , so folgt aus (47)

$$(50) \quad \begin{aligned} R(\bar{A}_i^{(n-h)}) &= \sigma_{ij}^{(h)} R(A_j^{(n-h)}) \\ &= \sigma_{ij}^{(h)} \varepsilon_{jp}^{(h+1)} A_p^{(n-h-1)} = \sigma_{ij}^{(h)} \varepsilon_{jp}^{(h+1)} \tau_{pq}^{(h+1)} \bar{A}_q^{(n-h-1)}, \end{aligned}$$

somit eine von (48), (49) verschiedene Randrelation.

Bezeichnen wir

$$(51) \quad \tilde{\varepsilon}_{iq}^{(h+1)} = \sigma_{ij}^{(h)} \varepsilon_{jp}^{(h+1)} \tau_{pq}^{(h+1)}, \text{ so lautet (50)}$$

$$(51') \quad R(A_i^{(n-h)}) = \tilde{\varepsilon}_{ij}^{(h+1)} \bar{A}_j^{(n-h-1)}$$

Fassen wir $\bar{A}_i^{(n-h)}$ auf als Basissystem von Σ^* und $\bar{\Sigma}^*$, so gelten für Σ^* die Randrelationen (51'), für $\bar{\Sigma}^*$ die Randrelationen (48).

Wenn wir aber im Sinne unserer abstrakten Topologie zwei Ringe Σ und $\bar{\Sigma}$ identisch nennen, wenn sich für Σ ein Basissystem $a_i^{(e)}$ und für $\bar{\Sigma}$ ein solches $\bar{a}_i^{(e)}$ so finden läßt, daß zugleich mit

$$R(a_j^{(e)}) = \varepsilon_{ij}^{(e)} a_j^{(e-1)} \text{ auch } R(\bar{a}_j^{(e)}) = \varepsilon_{ij}^{(e)} \bar{a}_j^{(e-1)}$$

gilt, so sind die Ringe Σ^* und $\bar{\Sigma}^*$ identisch.

Wir bezeichnen $\tilde{\varepsilon}_{ij}^{(h+1)} = \tilde{\varepsilon}_{ji}^{*(n-h)}$, dann lauten die Randrelationen (51') des Ringes Σ^*

$$(52) \quad \bar{R}(\bar{A}_i^{(n-h)}) = \tilde{\varepsilon}_{ji}^{*(n-h)} \bar{A}_j^{(n-h-1)}, \text{ wobei (51)}$$

$$(52') \quad \tilde{\varepsilon}_{qi}^{*(n-h)} = \sigma_{ij}^{*(n-h)} \varepsilon_{pj}^{*(n-h)} \tau_{pq}^{*(n-h-1)} \text{ gilt.}$$

Desgleichen gilt für den Ring Σ^* (48):

$$(53) \quad R(\bar{A}_i^{(n-h)}) = \bar{\varepsilon}_{ji}^{*(n-h)} \bar{A}_j^{(n-h-1)}, \text{ wo}$$

$$(53') \quad \bar{\varepsilon}_{qp}^{*(n-h)} = \sigma_{qr}^{*(n-h-1)} \varepsilon_{rs}^{*(n-h)} \tau_{sp}^{*(n-h)}, \text{ ist.}$$

Wegen $\tau_{sp}^{*(n-h)} \sigma_{pi}^{*(n-h)} = \delta_{si}$ folgt aus (53') $\sigma_{pi}^{*(n-h)} \bar{\varepsilon}_{qp}^{*(n-h)} =$
 $= \sigma_{qr}^{*(n-h-1)} \bar{\varepsilon}_{ri}^{*(n-h)}$ und weiter wegen $\sigma_{qr}^{*(n-h-1)} \tau_{jq}^{*(n-h-1)} = \delta_{rj}$,

$$(54) \quad \bar{\varepsilon}_{ji}^{*(n-h)} = \sigma_{pi}^{*(n-h)} \bar{\varepsilon}_{qp}^{*(n-h)} \tau_{jq}^{*(n-h-1)}.$$

Dies wird in (52') eingesetzt und hat

$$(54') \quad \bar{\varepsilon}_{qi}^{*(n-h)} = \sigma_{ij}^{*(n-h)} \sigma_{nj}^{*(n-h)} \bar{\varepsilon}_{vu}^{*(n-h)} \tau_{pv}^{*(n-h-1)} \tau_{pq}^{*(n-h-1)}.$$

Nun ist $\sigma_{ij}^{*(n-h)} \sigma_{uj}^{*(n-h)} = \alpha_{iu}^{(n-h)}$ eine unimodulare Matrix und
 $\tau_{pv}^{*(n-h)} \tau_{pq}^{*(n-h)} = \beta_{vq}^{(n-h)}$ ihre inverse. Man schreibt also (54')

$$(55) \quad \bar{\varepsilon}_{qi}^{*(n-h)} = \alpha_{iu}^{(n-h)} \bar{\varepsilon}_{vu}^{*(n-h)} \beta_{vq}^{(n-h-1)}, \quad i, u = 1, \dots, \alpha_n^* ;$$

$$v, q = 1, \dots, \alpha_{n-h-1}^*.$$

Führen wir im Ringe $\bar{\Sigma}^*$ die neue Basis

$$(56) \quad \tilde{A}_i^{(n-h)} = \alpha_{iu} \bar{A}_i^{(n-h)} \quad \text{ein, so gilt für diese Basis nach (8)}$$

$$(57) \quad R(\tilde{A}_i^{(n-h)}) = \bar{\varepsilon}_{ji}^{*(n-h)} \tilde{A}_j^{(n-h-1)},$$

woraus man wegen (52) auf die Identität der Ringe Σ^* und $\bar{\Sigma}^*$ schließt.

Neben dem dualen Ring ist für viele Probleme der μ dimensionale Sternring eines ν dimensionalen Komplexes $a_j^{(\nu)}$, wobei $\nu < \mu < n$ ist, von Bedeutung. Ehe wir ihn definieren, sei eine Erklärung für die Bedeutung der Aussage, „ $a_j^{(\nu)}$ liegt in $a_h^{(\mu)}$ “, oder was dem gleichbedeutend ist „ $a_h^{(\mu)}$ enthält $a_j^{(\nu)}$ “ gegeben.

Wir sagen: $a_h^{(\mu)}$ enthält $a_j^{(\nu)}$, wenn ein dem Rande von $a_h^{(\mu)}$ angehöriger, also $\mu - 1$ dimensionaler Komplex $a_j^{(\nu)}$ enthält.

Wir vervollständigen diese Erklärung: Gehört $a_j^{(\nu)}$ dem Rande eines $\nu + 1$ dimensional Komplexes $a_k^{(\nu+1)}$ an, so sagen wir, daß $a_k^{(\nu+1)}$ ihn enthalte.

Betrachten wir dann den μ dimensionalen Gerüststring eines n dimensional Ringes Σ und in ihm das Basiselement $a_j^{(\nu)}$, so gibt es ein bestimmtes Teilsystem der μ dimensional Basiselemente dieses Gerüstes, dessen Elemente alle $a_j^{(\nu)}$ enthalten.

Punkte verbindet, die diesen beiden Dreiecken entsprechen (z. B. die diese Punkte verbindende Strecke).

Einer Seite aber, die Rand nur eines Dreieckes des Komplexringes ist, ordnen wir eine Strecke zu, die einseitig berandet ist durch

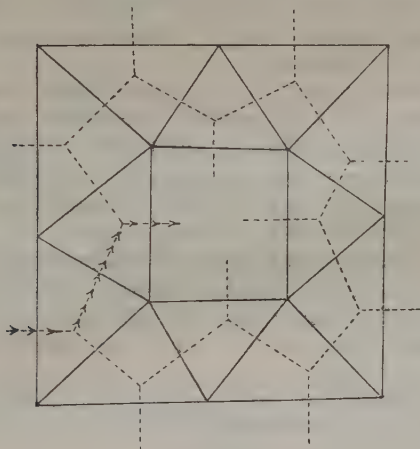


Fig. 2.

den diesem Dreieck entsprechenden Punkt. Einem Punkte des Komplexes wieder sind jene zweidimens. Elemente zugeordnet, die berandet sind durch alle jene Strecken, die den Strecken zugeordnet sind, die den betreffenden Punkt enthalten.

Enthält dabei ein ρ dimens. Element $M^{(\rho)}$ des Ringes Σ das $\rho-1$ dimens. $M^{(\rho-1)}$ 0, +1 oder -1mal, so soll das entsprechende $n-\rho$ dimens. Element $M^{*(n-\rho)}$ des dualen Ringes im entsprechenden $n-\rho+1$ dimens. Element $M^{*(n-\rho+1)}$, 0, +1 oder -1mal enthalten sein.

Wir haben in Fig. 1 und Fig. 2 die dualen Komplexringe gestrichelt gezeichnet. Die durch Pfeile markierten Kurvenzüge stellen eindimens. Zyklen dar, da ja die Endpunkte dieser Kurvenzüge den dualen Ringen nicht angehören.

Die Bettischen Zahlen des dualen Ringes des Ringes der Fig. 1

sind $h^*_0=0$, $h^*_1=0$, $h^*_2=1$ und die entsprechenden der Fig. 2

sind $h^*_0=0$, $h^*_1=1$, $h^*_2=1$.

Wir sagen, daß eine (z. B. in einem euklidischen R_N liegende) n dimens. Mannigfaltigkeit M_n ($n \leq N$) durch einen n dimens. geometrischen Komplexring (z. B. durch einen simplizialen Ring) approximiert ist, wenn folgende Bedingungen erfüllt sind:

1. Jeder Komplex des Ringes gehört der Mannigfaltigkeit an.

2. Jede unberandete ρ dimens. Mannigfaltigkeit F_ρ dieser M_n , $\rho < n$, ist homolog einem ρ dimens. Zyklus des Komplexringes.

Wir nennen hier zwei unberandete p dimens. Mannigfaltigkeiten homolog, wenn es in der M_n eine $p+1$ dimens. Mannigfaltigkeit gibt, deren Rand die Differenz dieser beiden Mannigfaltigkeiten ist.

Die p dimens. Mannigfaltigkeiten lassen sich in Klassen so zusammenfassen, daß die Elemente einer Klasse unter sich homolog sind.

Sind die Bedingungen 1) und 2) erfüllt, so entsprechen dann diese Klassen den Klassen der p^{ten} Homologiegruppe des approximierenden Ringes eineindeutig.

Der Komplexring der Fig. 1 approximiert dann jede abgeschlossene Kreisfläche, die ihn enthält, sein dualer dagegen die offene.

Der Komplexring der Fig. 2 approximiert die Fläche eines abgeschlossenen Kreisringes, sein dualer die des offenen.

Verallgemeinern wir unsere Überlegungen für n dimens. geometrische Komplexe, so erhalten wir das Resultat:

Ist M_n eine durch geometrische Komplexringe approximierbare Mannigfaltigkeit, M_i die Menge ihrer inneren Punkte, M_r die Menge ihrer Randpunkte, die Punkte von M_n sind, und M_e die Menge der übrigen Randpunkte, so stellt $M_i + M_r$ die duale Mannigfaltigkeit unserer $M_n = M_i + M_r$ dar.

Ist M_n unberandet, so ist $M_r = M_e = 0$ und M_n fällt mit seiner dualen Mannigfaltigkeit zusammen. (Letzteres Resultat stammt von Poincaré.)

III. ABSCHNITT.

Operationen an Ringen, die die Homologiegruppen unverändert lassen.

Es liege ein n dimensionaler Ring Σ vor, Σ_1 sei ein n dimensionaler echter Teilring von Σ , $\Sigma_1 < \Sigma$, dessen Basissystem Teilsystem ist eines bestimmten Basissystemes von Σ . Der Randring Σ_r von Σ_1 sei $n-1$ dimensional und der Komplementring Σ_e von Σ_1 in Σ demnach n dimensional.

Von Σ_1 und Σ_r setzen wir voraus^{*)}:

1) Σ_r ist ein Nullring, d. h. jeder p dimens. Zyklus von Σ_r für $0 < p < n-1$ sei in Σ_r homolog Null.

Ferner sei jeder nulldimens. Zyklus von Σ_r , der in Σ homolog Null ist, in Σ_r homolog Null.

2) Σ_1 ist ein Nullring, d. h. jeder p dimens. Zyklus von Σ_1 für $0 < p < n$ sei in Σ_1 homolog Null.

3) Jedem nulldimens. Zyklus von Σ_1 entspreche ein solcher in Σ_r , der mit ihm in Σ_1 homolog Null sei.

Die $n-1$ dimens. Zyklen des Ringes Σ_r haben eine Basis

$$(58) \quad C_{r1}^{(n-1)}, \quad C_{r2}^{(n-1)}, \quad \dots, \quad C_{ra}^{(n-1)}.$$

^{*)} Für $n=1$ fallen Voraussetzungen 1) und 2) weg.

deren Elemente als $n-1$ dimens. Zyklen von Σ_1 nach Voraussetzung 2) in Σ_1 homolog Null sind^{*)}:

$$(59) \quad C_{r_i}^{(n-1)} = R(A_{1i}^{(n)}), \quad i=1, \dots, \alpha.$$

Adjungieren wir zum Basissystem von Σ_2 die α Komplexe

$$(60) \quad A_{11}^{(n)}, \quad A_{12}^{(n)}, \dots, \quad A_{1\alpha}^{(n)}$$

von Σ_1 , so erhalten wir ein Basissystem für einen neuen Ring Σ_2 , dessen Homologiegruppen (also auch Bettische und Torsionszahlen) für $p < n$ mit denen des Ringes Σ identisch sind.

Weitere Elemente als die von (60) braucht man nicht zu adjungieren, da der Rand eines Elementes von (60) nach (59) bereits in Σ_2 , also in Σ_2 auftritt. (Die $n-1$ dimens. Gerüste von Σ_2 und Σ_3 sind identisch.)

Die Elemente (60) sind linear unabhängig, da aus $\pi_i A_{1i}^{(n)} = 0^{(n)}$, $i=1, \dots, \alpha$, durch Randbildung $\pi_i C_{r_i}^{(n-1)} = 0^{(n-1)}$, $i=1, \dots, \alpha$, folgt, d. h. $\pi_i = 0$, $i=1, \dots, \alpha$, infolge der Baseigenschaft der Elemente von (60).

Jeder p dimensionale Zyklus von Σ_2 ist wegen $\Sigma_2 < \Sigma$ p dimens. Zyklus von Σ , ist er in Σ_2 homolog Null, so ist er auch in Σ homolog Null.

Jede Klasse p dimens. Zyklen von Σ_2 ist demnach Teil einer Klasse der p dimens. Zyklen von Σ .

Wir lassen nun umgekehrt für $0 < p < n$ jedem p dimens. Zyklus von Σ einen solchen von Σ_2 entsprechen.

Sei $C^{(p)}$ dieser p dimens. Zyklus von Σ , so sei mit $C^{(p)}$ der ihm entsprechende von Σ_2 bezeichnet, zu dem wir auf folgendem Wege gelangen:

Indem wir $C^{(p)}$ durch die Basis des Ringes Σ ausdrücken, unterscheiden wir die drei Teilkomplexe von $C^{(p)}$

$$A_1^{(p)} \text{ in } \Sigma_1 = \Sigma_r, \quad A_2^{(p)} \text{ in } \Sigma_2 = \Sigma_r \text{ und } A_r^{(p)} \text{ in } \Sigma_r.$$

$A_r^{(p)}$ spalten wir seinerseits beliebig in zwei in Σ_r liegende Komplexe $A_{r1}^{(p)}$ und $A_{r2}^{(p)}$: $A_r^{(p)} = A_{r1}^{(p)} + A_{r2}^{(p)}$ und haben dann eine Zerlegung von $C^{(p)}$

^{*)} Für $n=1$: Die nulldimens. Zyklen von Σ_1 , die homolog Null in Σ_1 sind, haben eine Basis $C_{r_i}^{(0)}$, $i=1, \dots, \alpha$, usw.

Dabei berufen wir uns auf ein Resultat, das wir am Ende dieses Abschnittes in einer Bemerkung bringen: Die p dimens. Zyklen eines Teilringes, die im Ringe homolog Null sind, besitzen eine Basis.

$$(61) \quad C^{(q)} = k_1^{(q)} + k_2^{(q)}, \text{ wobei}$$

$$k_1^{(q)} = A_1^{(q)} + A_{r_1}^{(q)} \text{ und } k_2^{(q)} = A_2^{(q)} + A_{r_2}^{(q)} \text{ ist.}$$

Aus (61) erhalten wir durch Randbildung:

$$(61') \quad 0^{(q-1)} = R(k_1^{(q)}) + R(k_2^{(q)}).$$

Der $\rho-1$ dimens. Zyklus $R(k_1^{(q)})$, der in Σ_1 und Σ_2 liegt, ist somit ein Zyklus von Σ_r , und da er in Σ homolog Null ist und $\rho-1 < n-1$, ist er nach Voraussetzung 1) in Σ_r homolog Null, d. h. es gibt in Σ_r einen ρ dimens. Komplex $k_r^{(q)}$, dessen Rand $R(k_1^{(q)})$ ist:

$$(62) \quad R(k_1^{(q)}) = R(k_r^{(q)}).$$

Also ist

$$(62') \quad C_1^{(q)} = k_1^{(q)} - k_r^{(q)}$$

ρ dimens. Zyklus in Σ_1 und wegen $0 < \rho < n$ ist nach Voraussetzung

2) $C_1^{(q)} \sim 0$ in Σ_1 also auch ~ 0 in Σ .

Als den $C^{(q)}$ in Σ_3 entsprechenden ρ dimens. Zyklus definieren wir:

$$(62'') \quad \bar{C}^{(q)} = k_2^{(q)} + k_r^{(q)},$$

der im $n-1$ dimens. Gerüste von Σ_2 liegt.

Aus (61), (62') und (62'') folgt:

$$(63) \quad C^{(q)} = C_1^{(q)} + \bar{C}^{(q)}.$$

Die Zerlegung (63) des Zyklus $C^{(q)}$ in einen von Σ_1 und einen von Σ_2 ist nicht eindeutig, demnach auch nicht die Zuordnung $C^{(q)} \rightarrow \bar{C}^{(q)}$.

Sei neben (63): $C^{(q)} = \dot{C}_1^{(q)} + \dot{\bar{C}}^{(q)}$ eine solche zweite Zerlegung, also $\dot{C}_1^{(q)}$ ein Zyklus von Σ_1 und $\dot{\bar{C}}^{(q)}$ ein solcher von Σ_2 .

Dann folgt:

$$\bar{C}^{(q)} - \dot{\bar{C}}^{(q)} = \dot{C}_1^{(q)} - C_1^{(q)},$$

wo links ein Zyklus von Σ_1 und rechts ein Zyklus von Σ_2 steht.

Somit ist

$$\bar{C}^{(q)} - \dot{\bar{C}}^{(q)} = C_r^{(q)}$$

ein ρ dimens. Zyklus von Σ_r . Ist $\rho < n-1$, so ist $C_r^{(q)} \sim 0$ in Σ_r (Voraussetzung 1) also $\bar{C}^{(q)} \sim \bar{\dot{C}}^{(q)}$ in Σ_2 .

Ist $\rho = n-1$, so läßt sich $C_r^{(n-1)}$ durch die Basis (58) linear darstellen:

$$C_r^{(n-1)} = p_i C_{ri}^{(n-1)} i = 1 \dots \alpha,$$

woraus $C_r^{(n-1)} = R(p_i A_{1i}^{(n)}) \sim 0$ in Σ_3 folgt.

Es gilt somit stets

$$C^{(q)} \sim \bar{\dot{C}}^{(q)} \text{ in } \Sigma_3.$$

Jedem ρ dimens. Zyklus $C^{(q)}$ ($0 < \rho < n$) von Σ sind somit ρ dimens. Zyklen einer bestimmten Zyklenklasse von Σ_3 zugeordnet.

Aus (63) wieder folgt

$$(64) \quad C^{(q)} \sim \bar{C}^{(q)} \text{ in } \Sigma, \text{ da } C_1^{(q)}$$

nach Voraussetzung 2) in Σ homolog Null ist.

Es wird also jedem ρ dimens. Zyklus von Σ ($0 < \rho < n$) die Klasse¹⁰⁾ ρ dimens. Zyklen von Σ_3 zugeordnet, deren Elemente ganz in der diesen Zyklus enthaltenden Klasse von Σ liegen.

Wir können daraus aber noch nicht auf ein eindeutiges Entsprechen der ρ dimens. Zyklenklassen ($0 < \rho < n$) von Σ und Σ_3 schließen, da es ja möglich wäre, daß zwei verschiedenen ρ dimens. Zyklen $C_1^{(q)}$ und $C_2^{(q)}$, die einer Klasse in Σ angehören,

$$C_1^{(q)} \sim C_2^{(q)} \text{ in } \Sigma,$$

ρ dimens. Zyklen $C_1^{(q)}$ und $\bar{C}_2^{(q)}$ von Σ_3 zugeordnet sind, die in Σ_3 verschiedenen Zyklenklassen angehören.

Für letztere beiden Zyklen müßte dann gelten

$$\bar{C}_1^{(q)} \sim \bar{C}_2^{(q)} \text{ in } \Sigma \text{ (nach 64) aber } \bar{C}_1^{(q)} \not\sim \bar{C}_2^{(q)} \text{ in } \Sigma_3.$$

Betrachten wir dann den ρ dimens. Zyklus $Z^{(q)} = \bar{C}_1^{(q)} - \bar{C}_2^{(q)}$, so gälte:

$$(65) \quad \bar{Z}^{(q)} \sim 0 \text{ in } \Sigma, \quad \bar{Z}^{(q)} \not\sim 0 \text{ in } \Sigma_3.$$

Dann gäbe es in Σ einen Komplex $k^{(q+1)}$, dessen Rand $\bar{Z}^{(q)}$ wäre.

¹⁰⁾ „Die Klasse“ im Sinne von „nur Elemente dieser Klasse“.

Diesen Komplex denken wir uns in zwei Teilkomplexe $k_1^{(q+1)}$ in Σ_1 und $k_2^{(q+1)}$ in Σ_2 zerlegt, so daß dann

$$\bar{Z}^{(q)} = R(k_1^{(q+1)}) + R(k_2^{(q+1)}) \text{ gelten müßte.}$$

$R(k_1^{(q+1)})$ kann kein Nullkomplex sein, da sonst $Z^{(q)} \sim 0$ in Σ_3 wäre.

$Z^{(q)} - R(k_2^{(q+1)}) = R(k_1^{(q+1)})$ wäre also ρ dimens. Zyklus von Σ_r , der aber in Σ_r nicht homolog Null sein dürfte, da dies wieder $Z^{(q)} \sim 0$ in Σ_3 zur Folge hätte.

Es müßte demnach (Voraussetzung 1) $\rho = n - 1$ sein.

Dann aber gälte $R(k_1^{(n)}) = p_i C_{ri}^{(n-1)}$, $i = 1, \dots, \alpha$, also $R(k_1^{(n)}) = R(p_i A_{1i}^{(n)}) \sim 0$ in Σ_3 , woraus wieder $Z^{(q)} \sim 0$ in Σ_3 folgen würde. (Widerspruch.)

Es entsprechen somit, für $0 < \rho < n$, die ρ dimens. Zyklensklassen von Σ und Σ_3 einander eineindeutig und damit auch die entsprechenden Homologiegruppen.

Diese Behauptung beweisen wir nun auch für $\rho = 0$, wobei zum ersten Male die Voraussetzung 3) in Kraft tritt.

Sei $C^{(0)}$ ein nulldimens. Zyklus von Σ (= nulldimens. Komplex), so zerlegen wir ihn in seine drei Teilkomplexe (nulldimens. Zyklen) $\bar{C}_1^{(0)}$ in $\Sigma_1 - \Sigma_r$, $\bar{C}_2^{(0)}$ in $\Sigma_2 - \Sigma_r$ und $\bar{C}_r^{(0)}$ in Σ_r .

$$C^{(0)} = \bar{C}_1^{(0)} + \bar{C}_2^{(0)} + \bar{C}_r^{(0)}.$$

Zerlegen wir $\bar{C}_r^{(0)}$ beliebig in zwei in Σ_r liegende Bestandteile: $\bar{C}_r^{(0)} = \bar{C}_{r1}^{(0)} + \bar{C}_{r2}^{(0)}$, so erhalten wir

$$(66) \quad C^{(0)} = C_1^{(0)} + C_2^{(0)}, \text{ wo } C_1^{(0)} = \bar{C}_1^{(0)} + \bar{C}_{r1}^{(0)} \text{ und } C_2^{(0)} = \bar{C}_2^{(0)} + \bar{C}_{r2}^{(0)} \text{ ist.}$$

Nach Voraussetzung 3) gibt es in Σ_r einen nulldimens. Zyklus $C_r^{(0)}$, der in Σ_1 mit $C_1^{(0)}$ homolog ist:

$$C_1^{(0)} \sim C_r^{(0)} \text{ in } \Sigma_1.$$

Wir ordnen dann $C^{(0)}$ den Zyklus $C_2^{(0)} + C_r^{(0)}$ von Σ_2 (also auch von Σ_3) zu.

Sei $\bar{C}_2^{(0)} + \bar{C}_r^{(0)}$ ein zweiter $C^{(0)}$ so zugeordneter Zyklus von Σ_2 . Dann gilt

$$\begin{aligned} C^{(0)} &= C_2^{(0)} + C_r^{(0)} + (C_1^{(0)} - C_r^{(0)}) = C_2^{(0)} + C_r^{(0)} + R(k_1^{(1)}) \text{ und ebenso} \\ C^{(0)} &= \bar{C}_2^{(0)} + \bar{C}_r^{(0)} + R(\bar{k}^{(1)}). \end{aligned}$$

Daraus aber folgt

$$(C_2^{(0)} + C_r^{(0)}) - (\bar{C}_2^{(0)} + \bar{C}_r^{(0)}) = R(k_1^{(1)} - \bar{k}_1^{(1)})$$

Links steht ein Zyklus von Σ_2 und rechts ein Zyklus von Σ_1 . Somit ist $R(k_1^{(1)} - \bar{k}_1^{(1)})$ ein nulldimens. Zyklus von Σ_r , und da er in Σ_1 homolog Null ist, so ist er (Voraussetz. 1) homolog Null in $\Sigma_r^{(1)}$. Somit ist

$$C_2^{(0)} + C_r^{(0)} \sim \bar{C}_2^{(0)} + \bar{C}_r^{(0)} \text{ in } \Sigma_r, \text{ also auch in } \Sigma_3.$$

Also sind jedem nulldimens. Zyklus $C^{(0)}$ von Σ nulldimens. Zyklen einer bestimmten Klasse von Σ_3 zugeordnet.

Was die Eineindeutigkeit der Zuordnung der nulldimens. Zyklenklassen von Σ und Σ_3 betrifft, so könnte es zwei Zyklen von Σ_2 geben, die in Σ , aber nicht in Σ_3 homolog Null wären [analog (65)].

Dann gäbe es einen nulldimens. Zyklus $Z^{(0)}$ in Σ_2 , der in Σ , aber nicht in Σ_2 homolog Null wäre.

Aus $Z^{(0)} = R(k^{(1)})$ folgt, sobald man $k^{(1)} = k_1^{(1)} + k_2^{(1)}$ setzt:

$Z^{(0)} = R(k_1^{(1)}) + R(k_2^{(1)})$. $R(k_1^{(1)})$ ist dann nulldimens. Zyklus von Σ_r und (Voraussetzung 1) homolog Null in Σ_r . Somit gilt $Z^{(0)} \sim 0$ in Σ_3 , d.h. einen Zyklus $Z^{(0)}$ der beschriebenen Eigenschaft kann es nicht geben¹²⁾.

Es entsprechen somit für $\rho < n$ die ρ dimens. Zyklenklassen und demnach die entsprechenden Homologiegruppen der Ringe Σ und Σ_3 einander eineindeutig.

Was die n^{te} Homologiegruppe betrifft, so sei

$$(67) \quad a^{(n)}, a_2^{(n)}, a_r^{(n)}$$

die Basis der n dimens. Komplexe von Σ_1 . Da die Elemente von (60) n dimens. Komplexe von Σ_1 sind, so gilt

$$(68) \quad A_{1i} = p_{ij} a_j^{(n)}, \quad i=1, \dots, \alpha, \quad j=1, \dots, \varepsilon.$$

Es sei nun vorausgesetzt:

4) Jeder n dimensionale Zyklus von Σ enthalte Basiselemente der Reihe (67) nur in der Zusammenfassung (68).

Dann ist jeder n dimens. Zyklus von Σ ein solcher von Σ_3 und auch die n^{ten} Homologiegruppen dieser Ringe sind identisch.

¹¹⁾ Hier ist $n > 1$ vorausgesetzt. Für $n = 1$ ist $R(k_1^{(1)} - \bar{k}_1^{(1)})$ ein nulldimens. Zyklus von $\Sigma_r \sim 0$ in Σ_1 , also durch die $C_r^{(0)}$, $i=1, \dots, \alpha$, darstellbar, somit $= R(p_i A_{1i}^{(1)})$, also ~ 0 in Σ_3 .

¹²⁾ Für $n=1$: $R(k_1^{(1)})$ ist nulldimens. Zyklus von $\Sigma_r \sim 0$ in Σ_1 , also ~ 0 in Σ_3 .

Sind die Voraussetzungen 1), 2), 3) und 4) für das n dimens. Gerüst Σ eines N dimens. Ringes $\bar{\Sigma}$ erfüllt und ersetzt man dieses n dimens. Gerüst durch den Ring Σ_3 , so entsteht aus $\bar{\Sigma}$ ein N dimens. Ring $\bar{\Sigma}_3$.

Denn ist $a_i^{(n+1)}$ ein $n+1$ dimens. Basiselement von $\bar{\Sigma}$, so ist $R(a_i^{(n+1)})$ ein n dimens. Zyklus von Σ (dem n dimens. Gerüst von $\bar{\Sigma}$), somit ein n dimens. Zyklus von Σ_3 , also in $\bar{\Sigma}_3$ enthalten.

Damit ist die Ringeigenschaft von $\bar{\Sigma}$ festgestellt.

Wir behaupten weiter die Identität aller Homologiegruppen der Ringe $\bar{\Sigma}$ und $\bar{\Sigma}_3$,

Zu beweisen ist hier nur die Identität der n^{ten} Homologiegruppen. Nun ist jeder n dimens. Zyklus von $\bar{\Sigma}$ n dimens. Zyklus von $\bar{\Sigma}_3$ und umgekehrt. Ist $C^{(n)} \sim 0$ in $\bar{\Sigma}_3$, so ist er, da $\Sigma_3 < \bar{\Sigma}$ sicher homolog Null in $\bar{\Sigma}$. Ist $C^{(n)} \sim 0$ in Σ , so gilt $C^{(n)} = R(k^{(n+1)})$, wo $k^{(n+1)}$ ein $n+1$ dimens. Komplex von Σ , also auch von $\bar{\Sigma}_3$ ist, woraus $C^{(n)} \sim 0$ in Σ_3 folgt, q. e. d.

Bemerkung: Wir erhielten den Ring Σ_3 durch Adjunktion der Elemente (60) zum Basissystem des Ringes Σ_2 .

Bei aufmerksamer Lektüre unserer Überlegungen aber überzeugt man sich, daß man aus der Reihe der Elemente (58) jene hätte streichen können, die bereits in Σ_2 homolog Null sind und demzufolge die entsprechenden Elemente der Reihe (60)¹³⁾. Um möglichst wenig n dimens. Komplexe aus (60) adjungieren zu müssen, wird man eine Basisdarstellung (58) der $n-1$ dimens. Zyklen von Σ_r suchen, für die möglichst viele Elemente bereits in Σ_2 homolog Null sind.

Nach unserem Hilfssatz I sind wir in der Lage, ein Basissystem der $n-1$ dimens. Zyklen von Σ_2 aufzustellen, das das Basissystem (58) enthält. Sei

$$(69) \quad C_{r_1}^{(n-1)}, C_{r_2}^{(n-1)}, \dots, C_{r_\alpha}^{(n-1)}, C_{\alpha+1}^{(n-1)}, \dots, C_\beta^{(n-1)}$$

dieses Basissystem.

Ferner sei

$$(70) \quad N_i^{(n-1)} = \gamma_{ji} C_j^{(n-1)} \quad (C_j^{(n-1)} = C_{rj}^{(n-1)} \text{ für } j = 1, \dots, \alpha), \\ i = 1, \dots, \varepsilon, j = 1, \dots, \beta,$$

ein Basissystem der $n-1$ dimens. Zyklen von Σ_2 homolog Null in Σ_2 . Jeder Zyklus von Σ_r , der in $\Sigma_2 \sim 0$ ist, hat dann die Darstellung:

$$(71) \quad C_r^{(i-1)} = \pi_i N_i^{(n-1)} = \pi_i \gamma_{ji} C_j^{(n-1)}, \quad i = 1, \dots, \varepsilon, j = 1, \dots, \beta.$$

¹³⁾ Mit Ausnahme für die n^{te} Homologiegruppe.

Da er in Σ_r liegt, muß

$$(71') \quad \pi_i \gamma_{ji} = 0 \text{ sein, für } i = 1, \dots, \varepsilon, j = \alpha + 1, \alpha + 2, \dots, \beta.$$

Ist σ der Rang der Matrix $\|\gamma_{ji}\|$, $i = 1, \dots, \varepsilon$, $j = \alpha + 1, \alpha + 2, \dots, \beta$, so gibt es $\gamma = \varepsilon - \sigma$ unabhängige ganzzahlige Lösungssysteme $\pi_1, \pi_2, \dots, \pi_\varepsilon$, die eine Basis $\pi_1^{(\tau)}, \pi_2^{(\tau)}, \dots, \pi_\varepsilon^{(\tau)}$, $\tau = 1, \dots, \gamma$, besitzen, so daß dann

$$(72) \quad Z_{r\tau}^{(n-1)} = \pi_i^{(\tau)} N_i^{(n-1)}, \quad i = 1, \dots, \varepsilon, \tau = 1, \dots, \gamma$$

eine Basis der Zyklen ist von Σ_r , die in Σ_2 homolog Null sind.

Für die Zyklen (72) gilt:

$$(72') \quad Z_{r\tau}^{(n-1)} = c_{j\tau} C_{rj}^{(n-1)}, \quad \tau = 1, \dots, \gamma, j = 1, \dots, \alpha,$$

welches System wir durch unimodulare Transformation der Elemente (72) und (58) in die kanonische Gestalt bringen:

$$(72'') \quad \bar{Z}_{r1}^{(n-1)} = c_1 \bar{C}_{r1}^{(n-1)}, \quad \bar{Z}_{r2}^{(n-1)} = c_2 \bar{C}_{r2}^{(n-1)}, \dots, \bar{Z}_{r\gamma}^{(n-1)} = c_\gamma \bar{C}_{r\gamma}^{(n-1)}.$$

Nun lassen wir aber an Stelle des Basissystemes (58) das Basissystem $\bar{C}_{rj}^{(n-1)}$, $j = 1, \dots, \alpha$, treten und in diesem streichen wir alle Elemente $C_{rj}^{(n-1)}$, für die $c_j = 1$ ist, denn sie sind bereits in Σ_2 homolog Null.

So können wir z. B., ohne die Homologiegruppen $\rho = 0, 1$ zu ändern, aus der Kugeloberfläche eine Kalotte herauschneiden.

IV. ABSCHNITT.

Es seien Σ_1 und Σ_2 echte n dimens. Teilringe von Σ , deren Basissysteme in einem bestimmten Basissystem von Σ enthalten sind, und weiter soll jeder ρ dimens. Basiskomplex dieses Systems von Σ entweder in Σ_1 oder in Σ_2 oder in Σ_1 und Σ_2 enthalten sein.

Mit Σ_3 bezeichnen wir den Durchschnitttring von Σ_1 und Σ_2 .

Dann bedeute die symbolische Gleichung:

$$(73) \quad \Sigma = (\Sigma_1 - \Sigma_3) + (\Sigma_2 - \Sigma_3) + \Sigma_3,$$

daß jeder ρ dimens. Komplex von Σ *eindeutig* spaltbar ist in drei Teilkomplexe, die in $\Sigma_1 - \Sigma_3$, $\Sigma_2 - \Sigma_3$ und in Σ_3 liegen.

Wir ordnen dann jedem ρ dimens. Zyklus von Σ einen $\rho - 1$ dimens. Zyklus von Σ_3 zu: Sei $C^{(\rho)}$ dieser ρ dimens. Zyklus von Σ ,

so zerlegen wir ihn nach (73) in seine drei Bestandteile: $K_1^{(e)}$ in $\Sigma_1 - \Sigma_3$, $K_2^{(e)}$ in $\Sigma_2 - \Sigma_3$ und $K_3^{(e)}$ in Σ_3

$$(74) \quad C^{(e)} = K_1^{(e)} + K_2^{(e)} + K_3^{(e)}.$$

$K_3^{(e)}$ spalten wir willkürlich in zwei in Σ_3 liegende Komplexe $K_3^{(e)} = K_{31}^{(e)} + K_{32}^{(e)}$.

Setzen wir dann $K_1^{(e)} + K_{31}^{(e)} = H_1^{(e)}$, $K_2^{(e)} + K_{32}^{(e)} = H_2^{(e)}$, so gilt

$$(74') \quad C^{(e)} = H_1^{(e)} + H_2^{(e)}, \text{ wo } H_1^{(e)} \text{ in } \Sigma_1 \text{ und } H_2^{(e)} \text{ in } \Sigma_2 \text{ liegt.}$$

Randbildung von (74') ergibt

$$(75) \quad C^{(e-1)} = R(H_1^{(e)}) + R(H_2^{(e)}).$$

Somit ist $C_3^{(e-1)} = R(H_1^{(e)}) = R(-H_2^{(e)})$ ein $\rho-1$ dimens. Zyklus von Σ_3 , der in Σ_1 und in Σ_2 homolog Null ist.

Wir können somit jedem ρ dimens. Zyklus $C^{(e)}$ von Σ einen $\rho-1$ dimens. Zyklus $C_3^{(e-1)}$ in Σ_3 zuordnen (für $\rho > 0$), der in Σ_1 und in Σ_2 homolog Null ist. Diese Zuordnung ist ebenso wie die Spaltung (74') von $C^{(e)}$ nicht eindeutig.

Sei $C^{(e)} = \bar{H}_1^{(e)} + \bar{H}_2^{(e)}$ eine zweite solche Spaltung, so folgt dann

$H_1^{(e)} - \bar{H}_1^{(e)} = \bar{H}_2^{(e)} - H_2^{(e)} = H_3^{(e)}$, wo $H_3^{(e)}$ ein ρ dimens. Komplex in Σ_3 ist.

Somit gilt $R(H_1^{(e)}) - R(\bar{H}_1^{(e)}) = R(H_3^{(e)})$, d. h. $C_3^{(e-1)} \sim C_3^{(e-1)}$ in Σ_3 .

Alle auf obige Weise einem Zyklus $C^{(e)}$ von Σ zugeordneten Zyklen $C_3^{(e-1)}$, $\bar{C}_3^{(e-1)}$, ... gehören einer bestimmten Zyklenklasse von Σ_3 an.

Diese Zyklenklasse hat die bemerkenswerte Eigenschaft, daß jeder ihrer Zyklen homolog Null ist in Σ_1 und in Σ_2 .

Die Gesamtheit aller dieser Klassen stellt eine Untergruppe der $\rho-1$ ten Homologiegruppe von Σ_3 , wir bezeichnen sie mit Γ_3 .

Ihr Einheitsselement ist die Klasse der $\rho-1$ dimens. Zyklen von Σ_3 homolog Null in Σ_3 .

Definition: Wir nennen einen ρ dimens. Zyklus von Σ einen A-Zyklus, wenn er sich als Summe darstellen läßt eines ρ dimens. Zyklus von Σ_1 und eines ρ dimens. Zyklus von Σ_2 .

Jeder ρ dimens. Zyklus von Σ homolog Null in Σ ist ein A-Zyklus.

Sei $C^{(q)}$ dieser Zyklus, dann gilt $C^{(q)} = R(K^{(q+1)}) = R(K_1^{(q+1)}) + R(K_2^{(q+1)})$, wo $K^{(q+1)} = K_1^{(q+1)} + K_2^{(q+1)}$ die Spaltung von $K^{(q+1)}$ in einen Komplex von Σ_1 und einen von Σ_2 darstellt. Nun ist $R(K_1^{(q+1)}) = C_1^{(q)}$ und $R(K_2^{(q+1)}) = C_2^{(q)}$.

(Wir können demnach die letzte Behauptung schärfer formulieren: Jeder ρ dimens. Zyklus von $\Sigma \sim 0$ in Σ ist als Summe darstellbar eines ρ dimens. Zyklus von $\Sigma_1 \sim 0$ in Σ_1 und eines ρ dimens. Zyklus von $\Sigma_2 \sim 0$ in Σ_2 .)

Daraus folgt: Enthält eine Zyklenklasse der ρ^{ten} Homologiegruppe von Σ einen A-Zyklus, so sind alle Elemente A-Zyklen. Eine solche Klasse heie eine A-Klasse.

Die Gesamtheit der A-Klassen der ρ^{ten} Homologiegruppe von Σ stellt eine Untergruppe dar dieser Gruppe, wir bezeichnen sie mit Γ .

Zerlegen wir die ρ^{te} Homologiegruppe von Σ mittels dieser Untergruppe, so sei mit F die entsprechende Faktorgruppe bezeichnet.

Jedem $\rho-1$ dimens. Zyklus $C_3^{(q-1)}$ von Σ_3 , der in Σ_1 und in Σ_2 homolog Null ist, entspricht stets ein ρ dimens. Zyklus $C^{(q)}$, dem er zugeordnet ist.

Denn nach Annahme gibt es in Σ_1 wenigstens einen ρ dimens. Komplex $K_1^{(q)}$ und in Σ_2 einen solchen $-K_2^{(q)}$, da

$$C_3^{(q-1)} = R(K_1^{(q)}) = R(-K_2^{(q)}) \text{ ist.}$$

Dann ist $K_1^{(q)} + K_2^{(q)} = C^{(q)}$ ein ρ dimens. Zyklus, dem die Zyklenklasse von Σ_3 , die $C_3^{(q-1)}$ enthlt, zugeordnet ist.

Aber auch die Zuordnung $C_3^{(q-1)} \longrightarrow C^{(q)}$ ist nicht eindeutig.

Seien $\bar{K}_1^{(q)}$ in Σ_1 und $-\bar{K}_2^{(q)}$ in Σ_2 zwei weitere ρ dimens. Zyklen, deren Rand $C_3^{(q-1)}$ ist. Setzen wir dann $\bar{K}_1^{(q)} + \bar{K}_2^{(q)} = \bar{C}^{(q)}$, so gilt ebenfalls die Zuordnung $C_3^{(q-1)} \longrightarrow C^{(q)}$.

Nun ist $C^{(q)} - \bar{C}^{(q)} = (K_1^{(q)} - \bar{K}_1^{(q)}) + (K_2^{(q)} - \bar{K}_2^{(q)})$.

Da aber $R(K_1^{(q)}) = R(\bar{K}_1^{(q)})$ und $-R(K_2^{(q)}) = -R(\bar{K}_2^{(q)})$ gleich $C_3^{(q-1)}$ ist, so ist $K_1^{(q)} - \bar{K}_1^{(q)}$ ein ρ dimens. Zyklus von Σ_1 und $K_2^{(q)} - \bar{K}_2^{(q)}$ ein solcher von Σ_2 .

Daher ist $C^{(q)} - \bar{C}^{(q)}$ ein A-Zyklus.

Somit entspricht jedem ρ dimens. Zyklus von Σ_3 , der in Σ_1 und in Σ_2 homolog Null ist, eine Gesamtheit von ρ dimens. Zyklen in Σ der Eigenschaft, daß die Differenz je zweier Zyklen dieser Gesamtheit ein A-Zyklus ist.

(D. h. diese Gesamtheit ist in einem bestimmten Elemente der Faktorgruppe F enthalten.)

Seien jetzt $C_3^{(q-1)}$ und $\bar{C}_3^{(q-1)}$ zwei Zyklen einer Klasse von Γ_3 , d. h. es seien $C_3^{(q-1)}$ und $\bar{C}_3^{(q-1)}$ $\rho-1$ dimens. Zyklen von Σ_3 homolog Null in Σ_1 und in Σ_2 und $\bar{C}_3^{(q-1)}$ sei in Σ_3 $C_3^{(q-1)}$ homolog.

$$\text{Dann gilt } C_3^{(q-1)} = R(K_1^{(q)}) = R(-K_2^{(q)})$$

$$\bar{C}_3^{(q-1)} = R(\bar{K}_1^{(q)}) = R(-\bar{K}_2^{(q)}) \text{ nach Annahme.}$$

Also folgt für die zugeordneten ρ dimens. Zyklen von Σ ($C^{(q)} = K_1^{(q)} + K_2^{(q)}$, $\bar{C}^{(q)} = \bar{K}_1^{(q)} + \bar{K}_2^{(q)}$):

$$C^{(q)} - \bar{C}^{(q)} = (K_1^{(q)} - \bar{K}_1^{(q)}) + (K_2^{(q)} - \bar{K}_2^{(q)}).$$

Da weiter $C_3^{(q-1)} \sim \bar{C}_3^{(q-1)}$ in Σ_3 vorausgesetzt wurde, so gilt

$$R(K_1^{(q)} - \bar{K}_1^{(q)}) = R(K_3^{(q)}) \text{ resp. } R(-K_2^{(q)} + \bar{K}_2^{(q)}) = R(K_3^{(q)}).$$

Somit folgt

$$C^{(q)} - \bar{C}^{(q)} = (K_1^{(q)} - \bar{K}_1^{(q)} - K_3^{(q)}) + (K_2^{(q)} - \bar{K}_2^{(q)} + K_3^{(q)}) = C_1^{(q)} + C_2^{(q)}.$$

Es entsprechen also den $\rho-1$ dimens. Zyklen einer Klasse von Γ_3 ρ dimens. Zyklen in Σ , die alle in einem Elemente der Faktorgruppe F enthalten sind.

Seien umgekehrt $C^{(q)}$ und $\bar{C}^{(q)}$ zwei dimens. Zyklen von Σ , die zu einem Elemente dieser Faktorgruppe F gehören.

Dann gilt nach Annahme $C^{(q)} - \bar{C}^{(q)} = C_1^{(q)} + C_2^{(q)}$, wo $C_1^{(q)}$ in Σ_1 und $C_2^{(q)}$ in Σ_2 liegt.

Nun sei $C^{(q)}$ und $\bar{C}^{(q)}$ in Bestandteile von Σ_1 und Σ_2 zerlegt:

$$C^{(q)} = K_1^{(q)} + K_2^{(q)}, \quad \bar{C}^{(q)} = \bar{K}_1^{(q)} + \bar{K}_2^{(q)}.$$

Dann gilt die Zuordnung

$$C^{(q)} \longrightarrow R(K_1^{(q)}) = R(-K_2^{(q)}), \quad \bar{C}^{(q)} \longrightarrow R(\bar{K}_1^{(q)}) = R(-\bar{K}_2^{(q)}).$$

Aus $(K_1^{(e)} - \overline{K}_1^{(e)}) + (K_2^{(e)} - \overline{K}_2^{(e)}) = C_1^{(e)} + C_2^{(e)}$ schließen wir

$$K_1^{(e)} - \overline{K}_1^{(e)} - C_1^{(e)} = -K_2^{(e)} + \overline{K}_2^{(e)} + C_2^{(e)} = K_3^{(e)}, \text{ wo } K_3^{(e)} \text{ ein } \rho \text{ dimens.}$$

Komplex von Σ_3 ist. Randbildung hat dann

$$R(K_1^{(e)}) - R(\overline{K}_1^{(e)}) = R(K_3^{(e)}) \text{ zur Folge, d. h.}$$

$$R(K_1^{(e)}) \sim R(\overline{K}_1^{(e)}) \text{ in } \Sigma_3.$$

Den ρ dimens. Zyklen eines Elementes der Faktorgruppe F entsprechen somit $\rho-1$ dimens. Zyklen von Σ_3 , die in einem Elemente der Gruppe Γ_3 liegen.

Aus den beiden letzten Resultaten folgt dann:

Die Elemente der Faktorgruppe F und der Gruppe Γ_3 entsprechen einander ein-eindeutig.

$$\text{Da noch aus } C^{(e)} \longrightarrow C_3^{(e-1)} \text{ und } \overline{C}^{(e)} \longrightarrow \overline{C}_3^{(e-1)}$$

$$C^{(e)} + \overline{C}^{(e)} \longrightarrow C_3^{(e-1)} + \overline{C}_3^{(e-1)} \text{ folgt, so gilt der Satz:}$$

Die Gruppen F und Γ_3 sind einstufig isomorph bezogen.

Wir wollen nun eine Relation herstellen zwischen den Betti'schen Zahlen der Ringe Σ_1 , Σ_2 , Σ_3 und Σ .

Wir haben bereits festgestellt, daß jeder ρ dimens. Zyklus in Σ homolog Null in Σ sich als Summe darstellen läßt eines ρ dimens. Zyklus in Σ_1 homolog Null in Σ_1 und eines ρ dimens. Zyklus in Σ_2 homolog Null in Σ_2 .

Ist demnach

$$(76) \quad N_1^{(e)}, N_2^{(e)}, \dots, N_{\mu_Q}^{(e)} \text{ eine Basis der } \rho \text{ dimens. Zyklen von} \\ \Sigma_1 \sim 0 \text{ in } \Sigma_1$$

und

$$(76') \quad M_1^{(e)}, M_2^{(e)}, \dots, M_{r_Q}^{(e)} \text{ eine Basis der } \rho \text{ dimens. Zyklen von} \\ \Sigma_2 \sim 0 \text{ in } \Sigma_2,$$

so ist jeder ρ dimens. Zyklus in Σ homolog Null in Σ durch die Größen der Reihen (76) und (76') linear darstellbar.

Weiter können wir nach Hilfsatz I eine Basis der ρ dimens. Zyklen von Σ_1 und eine solche der ρ dimens. Zyklen von Σ_2 bilden, die ganz eine Basis der ρ dimens. Zyklen von Σ_3 enthalten.

Es sei

$$(77) \quad C_1^{(e)}, C_2^{(e)}, \dots, C_{\beta_Q}^{(e)} \text{ diese Basis der } \rho \text{ dimens. Zyklen von } \Sigma_3,$$

(77') $C_1^{(e)}, \dots, C_{\beta_e}^{(e)}, C_{\beta_e+1}^{(e)}, \dots, C_{\alpha_e}^{(e)}$ die betreffende der ρ dimens. Zyklen von Σ_1 und

(77'') $\bar{C}_1^{(e)}, \dots, \bar{C}_{\beta_e}^{(e)}, \bar{C}_{\beta_e+1}^{(e)}, \dots, \bar{C}_{\varepsilon_e}^{(e)}$ die betreffende der ρ dimens. Zyklen von Σ_2 , wo $\bar{C}_k^{(e)} = C_k^{(e)}$ für $k = 1, \dots, \beta_e$ ist.

Dann gilt

$$(78) \quad \begin{cases} N_j^{(e)} = a_{jk}^{(e)} C_k^{(e)}, j = 1, \dots, u_e, k = 1, \dots, \alpha_e \\ M_h^{(e)} = b_{hl}^{(e)} \bar{C}_l^{(e)}, h = 1, \dots, v_e, l = 1, \dots, \varepsilon_e. \end{cases}$$

Die ρ dimens. Zyklen $N_j^{(e)}$ und $M_h^{(e)}$ sind nun im allgemeinen nicht linear unabhängig.

Aus

$$(79) \quad p_j N_j^{(e)} + q_h M_h^{(e)} = 0^{(e)}, j = 1, \dots, u_e, h = 1, \dots, v_e,$$

folgt nach (78), da die Zyklen

$$(80) \quad C_1^{(e)}, C_2^{(e)}, \dots, C_{\beta_e}^{(e)}, C_{\beta_e+1}^{(e)}, \dots, C_{\alpha_e}^{(e)}, \bar{C}_{\beta_e+1}^{(e)}, \dots, \bar{C}_{\varepsilon_e}^{(e)}$$

linear unabhängig sind¹⁴⁾, das System:

$$(81) \quad \begin{cases} p_j a_{jk}^{(e)} + q_h b_{hk}^{(e)} = 0, k = 1, \dots, \beta_e \\ p_j a_{js}^{(e)} = 0, s = \beta_e + 1, \dots, \alpha_e, \\ q_h b_{ht}^{(e)} = 0, t = \beta_e + 1, \dots, \varepsilon_e. \\ j = 1, \dots, u_e, h = 1, \dots, v_e. \end{cases}$$

¹⁴⁾ Lineare Abhängigkeit der Elemente von (80) hätte $\sum_{h=1}^{\alpha_e} a_{hk} C_h^{(e)} = \sum_{t=\beta_e+1}^{\varepsilon_e} b_{kt} \bar{C}_t^{(e)}$ zur Folge. Aus dieser Relation folgt, daß der Zyklus $\left\{ \begin{array}{l} \text{links} \\ \text{rechts} \end{array} \right.$ in Σ_1 und Σ_2 , also in Σ_3 liegen müßte. Demnach gälte $\sum_{t=\beta_e+1}^{\varepsilon_e} b_{kt} \bar{C}_t^{(e)} = \sum_{j=1}^{\beta_e} c_j \bar{C}_j^{(e)}$, was, da (77'') eine Basis ist, $b_t = 0$, $t = \beta_e + 1, \dots, \varepsilon_e$, zur Folge hat. Damit aber gilt $\sum_{h=1}^{\alpha_e} a_{hk} C_h^{(e)} = 0^{(e)}$, was $a_h = 0$, $h = 1, \dots, \alpha_e$ zur Folge hat.

Sei σ_e der Rang des Systemes (81), $\sigma_e \leq u_e + v_e$, so gibt es genau $u_e + v_e - \sigma_e$ unabhängige Relationen (79), d. h. von den $u_e + v_e$ Zyklen $N_j^{(e)}$, $M_h^{(e)}$ sind $(u_e + v_e) - (u_e + v_e - \sigma_e) = \sigma_e$ linear unabhängig.

Wir erhalten auch sofort eine Basis der ρ dimens. Zyklen in Σ homolog Null in Σ , wenn wir durch unimodulare Transformationen der Elemente (80) und der $N_j^{(e)}$ und $M_h^{(e)}$ das System (78) in kanonische Form bringen.

Bezeichnen wir mit $\dot{C}_1^{(e)}, \dot{C}_2^{(e)}, \dots, \dot{C}_{\alpha_e + \epsilon_e - \beta_e}^{(e)}$ das neue Basissystem (80) und mit $\dot{N}_1^{(e)}, \dot{N}_2^{(e)}, \dots, \dot{N}_{u_e + v_e}^{(e)}$ das unimodular transformierte System der $N_j^{(e)}$ und $M_h^{(e)}$, so tritt an Stelle (78):

$$(82) \quad \dot{N}_i^{(e)} = c_i \dot{C}_i^{(e)}, \quad i = 1, \dots, \sigma_e, \quad c_i \neq 0,$$

wogegen $\dot{N}_{\sigma_e + \tau}^{(e)} = 0^{(e)}$ ist, für $\tau = 1, 2, \dots$

Da sich aber jeder ρ dimens. Zyklus von Σ homolog Null in Σ durch die $\dot{N}_i^{(e)}$, $i = 1, \dots, u_e + v_e$, ausdrücken läßt und $\dot{N}_{\sigma_e + \tau}^{(e)} = 0^{(e)}$ ist, so ist

$$(83) \quad \dot{N}_1^{(e)}, \quad \dot{N}_2^{(e)}, \quad \dot{N}_{\sigma_e}^{(e)}$$

eine Basis der ρ dimens. Zyklen von Σ homolog Null in Σ .

Somit ist σ_e die Basiszahl der ρ dimens. Zyklen in Σ homolog Null in Σ .

Wir schreiben die Relation (79) in der Form:

$$(84) \quad \Gamma^{(e)} = p_j N_j^{(e)} = -q_h M_h^{(e)}, \quad j = 1, \dots, u_e, \quad h = 1, \dots, v_e.$$

$\Gamma^{(e)}$ ist der allgemeinste ρ dimens. Zyklus, der in Σ_1 und Σ_2 , also in Σ_3 liegt und in Σ_1 und Σ_2 homolog Null ist (d. h. das allgemeinste Element einer Klasse von Γ_3).

Aus (81) gewinnt man ein Basissystem aller ganzzahligen Lösungssysteme von (81):

$$(85) \quad (t)p_1, (t)p_2, \dots, (t)p_{u_e}, (t)q_1, (t)q_2, \dots, (t)q_{v_e}$$

$t = 1, 2, \dots, u_e + v_e - \sigma_e$ und entsprechend ist

$$(86) \quad (t)\Gamma^{(e)} = (t)p_j N_j^{(e)} = - (t)q_h M_h^{(e)}, \quad j = 1, \dots, u_e, \quad h = 1, \dots, v_e$$

eine Basis der ρ dimens. Zyklen von Σ_3 , die in Σ_1 und Σ_2 homolog Null sind.

Bezeichnen wir mit τ_ρ die Basiszahl der ρ dimens. Zyklen von Σ_3 homolog Null in Σ_1 und Σ_2 , so gilt:

$$(87) \quad \tau_\rho = u_\rho + v_\rho - \sigma_\rho.$$

Die Differenz: Basiszahl der ρ dimens. Zyklen von $\Sigma_1 \sim 0$ in Σ_1 plus Basiszahl der ρ dimens. Zyklen von $\Sigma_2 \sim 0$ in Σ_2 minus Basiszahl der ρ dimens. Zyklen von $\Sigma \sim 0$ in Σ ist gleich der Basiszahl der ρ dimens. Zyklen von Σ_3 homolog Null in Σ_1 und Σ_2 .

Es sei nun

$$(86') \quad (t)\Gamma^{(\rho-1)}, \quad t = 1, \dots, \tau_{\rho-1}$$

eine Basis der $\rho-1$ dimens. Zyklen von Σ_3 homolog Null in Σ_1 und Σ_2 und

$$(86'') \quad (s)P^{(\rho-1)}, \quad s = 1, \dots, w_{\rho-1}$$

eine Basis der $\rho-1$ dimens. Zyklen von Σ_3 homolog Null in Σ_3 .

Dann gilt:

$$(88) \quad (s)P^{(\rho-1)} = d_{st}^{(\rho-1)} (t)\Gamma^{(\rho-1)}, \quad s = 1, \dots, w_{\rho-1}, \quad t = 1, \dots, \tau_{\rho-1},$$

denn jeder $\rho-1$ dimens. Zyklus ~ 0 in Σ_3 ist natürlich auch ~ 0 in Σ_1 und Σ_2 .

Das System (88) bringen wir durch unimodulare Transformation der Elemente (86') und (86'') in die kanonische Form:

$$(88') \quad (1)\dot{P}^{(\rho-1)} = d_{11} (1)\dot{\Gamma}^{(\rho-1)}, \dots, (w_{\rho-1})\dot{P}^{(\rho-1)} = d_{w_{\rho-1} w_{\rho-1}} (w_{\rho-1})\dot{\Gamma}^{(\rho-1)}$$

(88') sind dann die definierenden Relationen der Elemente der Gruppe Γ_3 in bezug auf die Basis

$$(89) \quad (t)\dot{\Gamma}^{(\rho-1)}, \quad t = 1, \dots, \tau_{\rho-1}.$$

Sei nun

$$(90) \quad (t)\Gamma^{(\rho)}, \quad t = 1, \dots, \tau_{\rho-1}$$

der dem $\rho-1$ dimens. Zyklus $(t)\dot{\Gamma}^{(\rho-1)} \sim 0$ in Σ_1 und Σ_2 zugeordnete ρ dimens. Zyklus von Σ .

Adjungiert man diese $\tau_{\rho-1}$ Elemente zu den $\alpha_\rho + \varepsilon_\rho - \beta_\rho$ Elementen der Reihe (80), so erhält man $\alpha_\rho + \varepsilon_\rho - \beta_\rho + \tau_{\rho-1}$ ρ dimens. Zyklen in Σ der Eigenschaft, daß jeder ρ dimens. Zyklus in Σ durch sie linear darstellbar ist.

Denn sei $C^{(q)}$ ein beliebiger ρ dimens. Zyklus in Σ und $C_3^{(q-1)}$ der ihm zugeordnete $\rho-1$ dimens. Zyklus in Σ_3 , der also in Σ_1 und Σ_2 homolog Null ist, so gilt

$$C_3^{(q-1)} = a_{t(t)} \dot{\Gamma}^{(q-1)}, \quad t = 1, \dots, \tau_{q-1}.$$

Dem Zyklus $a_{t(t)} \dot{\Gamma}^{(q-1)}$ sind dann in Σ der Zyklus $C^{(q)}$ und auch der Zyklus $a_{t(t)} \Gamma^{(q)}$ zugeordnet. Somit ist

$$C^{(q)} - a_{t(t)} \Gamma^{(q)}, \quad t = 1, \dots, \tau_{q-1}$$

ein A-Zyklus und somit durch die Elemente der Reihe (80) linear darstellbar. Aber diese $\alpha_q + \varepsilon_q - \beta_q + \tau_{q-1}$ ρ dimens. Zyklen sind nicht linear unabhängig. Denn aus (88') folgt, daß

$$d_{s(s)} \Gamma^{(q)}, \quad s = 1, \dots, w_{q-1}, \quad \text{über } s \text{ nicht summiert}$$

für jedes $s = 1, \dots, w_{q-1}$, ein A-Zyklus ist. Daher gilt:

$$(91) \quad d_{s(s)} \Gamma^{(q)} = b_{sh} C_h^{(q)}, \quad h = 1, 2, \dots, \alpha_q + \varepsilon_q - \beta_q, \quad s = 1, \dots, w_{q-1}$$

[in (91) haben wir die Zyklen der Reihe (80) fortlaufend numeriert].

Umgekehrt folgt aus $\sum_{t=1}^{\tau_{q-1}} p_{t(t)} \Gamma^{(q)} = \sum_{h=1}^{\alpha_q + \varepsilon_q - \beta_q} c_h C_h^{(q)}$, die Relation $\sum_{t=1}^{\tau_{q-1}} p_{t(t)} \dot{\Gamma}^{(q-1)} \sim 0$ in Σ_3 , demnach nach (88')

$$p_t = q_t d_t \quad (\text{über } t \text{ nicht summiert}), \quad t = 1, \dots, w_{q-1},$$

$$p_t = 0, \quad t > w_{q-1}.$$

Somit sind (91) die einzigen Beziehungen zwischen $\alpha_q + \varepsilon_q - \beta_q + \tau_{q-1}$ ρ dimens. Zyklen und sie sind [man beachte die Matrix (91)] linear unabhängig. Bezeichnen wir nun fortlaufend numerierend mit $C_h^{(q)}$, $h = 1, 2, \dots, \alpha_q + \varepsilon_q - \beta_q + \tau_{q-1}$ diese ρ dimens. Zyklen, so hat das System (91) die Gestalt:

$$(92) \quad c_{hs} C_h^{(q)} = 0^{(q)}, \quad s = 1, \dots, w_{q-1}, \quad h = 1, \dots, \alpha_q + \varepsilon_q - \beta_q + \tau_{q-1},$$

Durch unimodulare Transformation der $C_h^{(q)}$ bringen wir (92) in die kanonische Form:

$$(93) \quad c_1 \dot{C}_1^{(q)} = 0^{(q)}, \quad c_2 \dot{C}_2^{(q)} = 0^{(q)}, \dots, \quad c_{w_{q-1}} \dot{C}_{w_{q-1}}^{(q)} = 0^{(q)},$$

woraus, da die c_i , $i = 1, \dots, w_{e-1}$, von Null verschieden sind:

$$(93') \quad \dot{C}_s^{(e)} = 0^{(e)}, \quad s = 1, \dots, w_{e-1} \text{ folgt.}$$

Somit ist

$$(94) \quad \dot{C}_h^{(e)}, \quad h = w_{e-1} + 1, \dots, \alpha_e + \varepsilon_e - \beta_e + \tau_{e-1}$$

ein System von ρ dimens. Zyklen von Σ , zwischen denen es keine lineare Beziehung mehr gibt, und durch das jeder ρ dimens. Zyklus linear darstellbar ist. (94) ist somit ein Basissystem der ρ dimens. Zyklen von Σ .

Da $\sigma_e = u_{(e)} + v_e - \tau_e$ die Basiszahl der ρ dimens. Zyklen von Σ homolog Null in Σ ist, so ist $h_{e(\Sigma)}$, die ρ^{te} Bettische Zahl von Σ

$$h_{e(\Sigma)} = \alpha_e + \varepsilon_e - \beta_e + \tau_{e-1} - w_{e-1} - u_{(e)} - v_e + \tau_e.$$

Mit γ_e bezeichnen wir die Differenz $\tau_e - w_e$, d. i. die Basiszahl der ρ dimens. Zyklen von $\Sigma_3 \infty 0$ in Σ_1 und in Σ_2 minus der Basiszahl der ρ dimens. Zyklen von $\Sigma_3 \infty 0$ in Σ_3 .

$$(95) \quad \gamma_e = \tau_e - w_e.$$

Dann gilt

$$h_{e(\Sigma)} = (\alpha_e - \mu_e) + (\varepsilon_e - v_e) - (\beta_e - w_e) + (\tau_e - w_e) + (\tau_{e-1} - w_{e-1}),$$

also

$$(96) \quad \boxed{h_{e(\Sigma)} = h_{e(\Sigma_1)} + h_{e(\Sigma_2)} - h_{e(\Sigma_3)} + \gamma_e + \gamma_{e-1}.}$$

In (96) bedeutet $h_{e(\Sigma_i)}$ die ρ^{te} Bettische Zahl des Ringes Σ_i , $i = 1, 2, 3$.

Bemerkung: Die Formel (96) gilt, wenn man $\gamma_{-1} = 0$ setzt, auch für $\rho = 0$.

Denn ist α_0 die Basiszahl der 0 dimens. Zyklen von Σ_1

ε_0	"	"	"	"	"	"	"	Σ_2
β_0	"	"	"	"	"	"	"	Σ_3 , so ist
$\alpha_0 + \varepsilon_0 - \beta_0$	"	"	"	"	"	"	"	Σ und
$\sigma_0 = \mu_0 + v_0 - \tau_0$	"	"	"	"	"	"	"	$\Sigma \infty 0$ in Σ .

Somit ist

$$\begin{aligned} h_{0(\Sigma)} &= (\alpha_0 - \mu_0) + (\varepsilon_0 - v_0) - (\beta_0 - w_0) + (\tau_0 - w_0) \\ &= h_{0(\Sigma_1)} + h_{0(\Sigma_2)} - h_{0(\Sigma_3)} + \gamma_0. \end{aligned}$$

Neben $\gamma_{-1} = 0$ ist auch $\gamma_n = 0$, da es ja im n dimens. Ringe Σ keine n dimens. Zyklen homolog Null gibt.

Beispiel: Zwei gebogene Röhrenflächen Σ_1 und Σ_2 seien laut Figur übereinandergestülpt. Sie bilden zusammen die Ringfläche Σ .

Σ_3 besteht hier aus zwei getrennten Röhrenstücken.

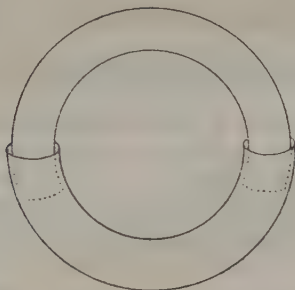


Fig. 3.

Es ist hier: $h_2(\Sigma) = 1$, $h_2(\Sigma_1) = h_2(\Sigma_2) = h_2(\Sigma_3) = \gamma_2 = 0$, $\gamma_1 = 1$

$h_1(\Sigma) = 2$, $h_1(\Sigma_1) = h_1(\Sigma_2) = 1$, $h_1(\Sigma_3) = 2$, $\gamma_1 = 1$, $\gamma_0 = 1$

$h_0(\Sigma) = 1$, $h_0(\Sigma_1) = h_0(\Sigma_2) = 1$, $h_0(\Sigma_3) = 2$, $\gamma_0 = 1$, $\gamma_{-1} = 0$.

Bemerkung: Die Relation (96) hätten wir auch folgendermaßen herleiten können. Es sei G eine kommutative Gruppe und A eine Untergruppe von G . Wir zerlegen G mittels A in Nebengruppen

$$G = (A + h_0) + (A + h_1) + (A + h_2) + \dots,$$

die als Elemente betrachtet eine Gruppe, die Faktorgruppe von G in bezug auf A , bilden.

Es sei nun A (d. h. die Elemente von A) darstellbar durch α Elemente

$$a_1, a_2, \dots, a_\alpha,$$

zwischen denen α_1 definierende Relationen

$$c_{pi} a_i = 0, \quad p = 1, \dots, \alpha_1, \quad \text{bestehen mögen.}$$

Desgleichen sei die Faktorgruppe durch β Elemente darstellbar

$$A + h_1, A + h_2, \dots, A + h_\beta,$$

mit β_1 definierenden Relationen

$$d_{qi}(A + b_i) = 0, \quad q = 1, \dots, \beta_1.$$

Wir behaupten dann, daß die Elemente von G durch $\alpha + \beta$ Elemente darstellbar sind, zwischen denen $\alpha_1 + \beta_1$ definierende Relationen bestehen.

Beweis: Irgend ein Element von G ist in einer Nebengruppe enthalten und hat daher die Form $\tilde{a} + \tilde{h}$. Für $A + \tilde{h}$ gilt die Darstellung

$$A + \tilde{h} = e_i(A + h_i) \quad i = 1, \dots, \beta$$

Nehmen wir in $A + h_i$ das beliebige Element $\bar{a}_i + h_i$ (z. B. $0 + h_i = h_i$), so ist $\sum_{i=1}^{\beta} e_i(\bar{a}_i + h_i)$ sicher in $A + \tilde{h}$ enthalten, also gleich $\tilde{a} + \tilde{h}$.

Ferner ist $(\tilde{a} + \tilde{h}) - (\bar{a} + \tilde{h}) = \tilde{a} - \bar{a} = \sum_{i=1}^{\alpha} \pi_i a_i$, also

$$\tilde{a} + \tilde{h} = \sum_{i=1}^{\alpha} \pi_i a_i + \sum_{i=1}^{\beta} e_i(\bar{a}_i + h_i).$$

Da aber $\bar{a}_i$ durch die a_i selbst linear darstellbar ist, so folgt hieraus, daß jedes Element von G durch die $\alpha + \beta$ Elemente a_i und h_i linear darstellbar ist. Zwischen diesen $\alpha + \beta$ Elementen bestehen als Relationen:

$$c_{pi} a_i = 0, \quad p = 1, \dots, \alpha_1$$

$$d_{qj}(\bar{a}_j + h_j) = \text{Nullelement, d. h. Element von } A = \tilde{f}_{qt} a_t, \text{ oder } d_{qj} h_j = \tilde{f}_{qt} a_t, \quad q = 1, \dots, \beta_1., \text{ q. e. d.}$$

In unserem Fall ist G die ρ^{te} Homologiegruppe von Σ , A die Gruppe Γ und F die Faktorgruppe.

$$\text{Daher ist } \alpha = \alpha_e + \varepsilon_e - \beta_e, \quad \alpha_1 = \sigma_e, \quad \beta = \tau_{e-1}, \quad \beta_1 = w_{e-1}.$$

$$\begin{aligned} \alpha + \beta - \alpha_1 - \beta_1 &= \alpha_e + \varepsilon_e - \beta_e - \sigma_e + \gamma_{e-1}, \quad (\tau_{e-1} - w_{e-1} = \gamma_{e-1}) \\ &= \alpha_e + \varepsilon_e - \beta_e - \mu_e - v_e + \tau_e + \gamma_{e-1} \\ &= (\alpha_e - \mu_e) + (\varepsilon_e - v_e) - (\beta_e - w_e) + (\tau_e - w_e) + \gamma_{e-1} \\ &= h_{e(\Sigma_1)} + h_{e(\Sigma_2)} - h_{e(\Sigma_3)} + \gamma_e + \gamma_{e-1}. \end{aligned}$$

(Eingegangen: 16. XI. 1927.)

Über eine Verallgemeinerung des Poissonschen Theorems.

Von E. Schuntner, Wien.

Die Theorie der Integralinvarianten gestattet nicht nur, das bekannte Theorem von Jacobi — wonach man ein drittes Integral eines kanonischen Systems gewöhnlicher Differentialgleichungen angeben kann, wenn man zwei Integrale kennt — in sehr einfacher Weise abzuleiten, sondern auch es nach einer gewissen Richtung hin zu verallgemeinern.

Ist das System gewöhnlicher Differentialgleichungen

$$\frac{dx_i}{X_i} = dt \quad i = 1, 2, \dots, n \quad (1)$$

gegeben, und versteht man unter $M_{r_1 r_2 \dots r_\lambda}$ Funktionen von $x_1 x_2 \dots x_n t$, so besteht die Bedingung dafür, daß

$$\int_{R\lambda} \Sigma M_{r_1 r_2 \dots r_\lambda} \delta x_{r_1} \delta x_{r_2} \dots \delta x_{r_\lambda}$$

eine absolute Integralinvariante λ^{ter} Ordnung darstellt (die Integration ist über einen beliebigen λ -fach ausgedehnten Raum zu erstrecken) in dem identischen Erfülltsein des Systems von Differentialgleichungen

$$\begin{aligned} \frac{d}{dt} M_{r_1 r_2 \dots r_\lambda} + \sum_{i=1}^n \left(M_{i r_2 r_3 \dots r_\lambda} \frac{\partial X_i}{\partial x_{r_1}} + M_{r_1 i r_3 \dots r_\lambda} \frac{\partial X_i}{\partial x_{r_2}} + \dots \right. \\ \left. \dots + M_{r_1 \dots r_{\lambda-1} i} \frac{\partial X_i}{\partial x_{r_\lambda}} \right) = 0 \end{aligned} \quad (2)$$

zusammen mit (1). Die Indizes $r_1 r_2 \dots r_\lambda$ durchlaufen dabei alle möglichen Kombinationen der n Elemente $1 2 \dots n$ zu je λ .

Ist das System (1) ein kanonisches, d. h. hat es die Form

$$\frac{\frac{dq_i}{\partial H}}{\frac{\partial p_i}{\partial q_i}} = \frac{\frac{dp_i}{\partial H}}{\frac{\partial H}{\partial q_i}} = dt \quad i = 1, 2 \dots m, \quad (3)$$

und betrachtet man Integralinvarianten der Ordnung $\lambda = 2k$ dieses Systems, wobei man den Koeffizienten von

$$\delta q_{\alpha_1} \delta q_{\alpha_2} \dots \delta q_{\alpha_v} \delta p_{\beta_1} \delta p_{\beta_2} \dots \delta p_{\beta_\mu}$$

$$v + \mu = 2k$$

mit

$$M [q_{\alpha_1} q_{\alpha_2} \dots q_{\alpha_v} p_{\beta_1} p_{\beta_2} \dots p_{\beta_\mu}]$$

bezeichnet, so erfüllen wegen (2) die Koeffizienten von $\delta q_{\alpha_1} \delta q_{\alpha_2} \dots \delta q_{\alpha_k}$ $\delta p_{\alpha_1} \delta p_{\alpha_2} \dots \delta p_{\alpha_k}$, d. h. die Funktionen $M [q_{\alpha_1} q_{\alpha_2} \dots q_{\alpha_k} p_{\alpha_1} p_{\alpha_2} \dots p_{\alpha_k}]$ der Argumente $t, q_1 \dots q_m p_1 \dots p_m$ identisch das System

$$\begin{aligned} & - \frac{d}{dt} M [q_{\alpha_1} q_{\alpha_2} \dots q_{\alpha_k} p_{\alpha_1} p_{\alpha_2} \dots p_{\alpha_k}] = \\ & \sum_{i=1}^m M [q_i q_{\alpha_2} \dots q_{\alpha_k} p_{\alpha_1} \dots p_{\alpha_k}] \frac{\partial^2 H}{\partial p_i \partial q_{\alpha_1}} + \\ & + \sum_{i=1}^m M [q_{\alpha_1} q_i q_{\alpha_3} \dots q_{\alpha_k} p_{\alpha_1} \dots p_{\alpha_k}] \frac{\partial^2 H}{\partial p_i \partial q_{\alpha_2}} + \dots \\ & \dots + \sum_{i=1}^m M [q_{\alpha_1} \dots q_{\alpha_{k-1}} q_i p_{\alpha_1} \dots p_{\alpha_k}] \frac{\partial^2 H}{\partial p_i \partial q_{\alpha_k}} + \\ & - \sum_{i=1}^m M [p_i q_{\alpha_2} \dots q_{\alpha_k} p_{\alpha_1} \dots p_{\alpha_k}] \frac{\partial^2 H}{\partial q_i \partial q_{\alpha_1}} - \\ & - \sum_{i=1}^m M [q_{\alpha_1} p_i q_{\alpha_3} \dots q_{\alpha_k} p_{\alpha_1} \dots p_{\alpha_k}] \frac{\partial^2 H}{\partial q_i \partial q_{\alpha_2}} - \dots \\ & \dots - \sum_{i=1}^m M [q_{\alpha_1} \dots q_{\alpha_{k-1}} p_i p_{\alpha_1} \dots p_{\alpha_k}] \frac{\partial^2 H}{\partial q_i \partial q_{\alpha_k}} \\ & + \sum_{i=1}^m M [q_{\alpha_1} \dots q_{\alpha_k} q_i p_{\alpha_2} \dots p_{\alpha_k}] \frac{\partial^2 H}{\partial p_i \partial p_{\alpha_1}} + \\ & + \sum_{i=1}^m M [q_{\alpha_1} \dots q_{\alpha_k} p_{\alpha_1} q_i p_{\alpha_3} \dots p_{\alpha_k}] \frac{\partial^2 H}{\partial p_i \partial p_{\alpha_2}} \\ & \dots + \sum_{i=1}^m M [q_{\alpha_1} \dots q_{\alpha_k} p_{\alpha_1} \dots p_{\alpha_{k-1}} q_i] \frac{\partial^2 H}{\partial p_i \partial p_{\alpha_k}} \\ & - \sum_{i=1}^m M [q_{\alpha_1} \dots q_{\alpha_k} p_i p_{\alpha_2} \dots p_{\alpha_k}] \frac{\partial^2 H}{\partial q_i \partial p_{\alpha_1}} - \\ & - \sum_{i=1}^m M [q_{\alpha_1} \dots q_{\alpha_k} p_{\alpha_1} p_i \dots p_{\alpha_k}] \frac{\partial^2 H}{\partial q_i \partial p_{\alpha_2}} \\ & \dots - \sum_{i=1}^m M [q_{\alpha_1} q_{\alpha_2} \dots q_{\alpha_k} p_{\alpha_1} \dots p_{\alpha_{k-1}} p_i] \frac{\partial^2 H}{\partial q_i \partial p_{\alpha_k}}, \end{aligned} \quad (4)$$

das mit (3) zusammen besteht. Das System besteht aus $\binom{m}{k}$ Gleichungen und die Funktionen M ändern das Zeichen, wenn man in der eckigen Klammer zwei Buchstaben vertauscht, die ja hier die Rolle von Indizes spielen. Macht man, nachdem man alle Gleichungen (4) addiert hat, von dieser Eigenschaft Gebrauch, so folgt

$$\frac{d}{dt} \Sigma M [q_{\alpha_1} q_{\alpha_2} \dots q_{\alpha_k} p_{\alpha_1} \dots p_{\alpha_k}] = 0.$$

In der Tat tilgen sich alle Glieder, die $\frac{\partial^2 H}{\partial q \partial q}$ als Faktor enthalten, gegenseitig, ferner alle mit $\frac{\partial^2 H}{\partial q \partial p}$ und alle mit $\frac{\partial^2 H}{\partial p \partial p}$. Setzt man noch kurz $M [q_{\alpha_1} q_{\alpha_2} \dots q_{\alpha_k} p_{\alpha_1} \dots p_{\alpha_k}] = M_{\alpha_1 \alpha_2 \dots \alpha_k}$, so kann man das Ergebnis folgendermaßen aussprechen:

Kennt man von einem kanonischen System (3) eine Integralinvariante gerader, etwa $2k^{\text{ter}}$ Ordnung, so stellt

$$\underline{\Sigma M_{\alpha_1 \alpha_2 \dots \alpha_k} = c}$$

ein Integral dieses Systems dar.

Es läßt sich unmittelbar erkennen, daß dieser Satz das Poissonsche Theorem in sich schließt. Denn bedeuten $\varphi_1 = c_1$ und $\varphi_2 = c_2$ zwei Integrale des kanonischen Systems (3), so kann man sofort eine Integralinvariante 2^{ter} Ordnung angeben. Sie lautet:

$$\int_{K_2} \sum_{i=1}^m \frac{\partial (\varphi_1 \varphi_2)}{\partial (q_i p_i)} \delta q_i \delta p_i,$$

die Koeffizienten M sind von der Form $M [q_i p_i] = M_i = \frac{\partial (\varphi_1 \varphi_2)}{\partial (q_i p_i)}$ und unser Theorem geht über in das folgende:

„Kennt man von einem kanonischen System (3) zwei unabhängige Integrale $\varphi_1 = c_1$ und $\varphi_2 = c_2$, so repräsentiert

$$\underline{(\varphi_1 \varphi_2) = \Sigma \frac{\partial (\varphi_1 \varphi_2)}{\partial (q_i p_i)} = c}$$

ein neues Integral von (3).“

Das ist der bekannte Poissonsche Satz, der hier als trivialer Spezialfall erscheint.

Diese Behauptungen haben ein Gegenstück in der Theorie der Variationslösungen, d. h. jener Funktionen $\xi_{r_1 r_2 \dots r_\lambda}$ von $x_1 x_2 \dots x_n t$, die, mit (1) zusammen, das System

$$\frac{d}{dt} \xi_{r_1 r_2 \dots r_\lambda} - \sum_{i=1}^m \left(\xi_{i r_2 \dots r_\lambda} \frac{\partial X_{r_1}}{\partial x_i} + \xi_{r_1 i r_3 \dots r_\lambda} \frac{\partial X_{r_2}}{\partial x_i} + \dots \right. \\ \left. \dots + \xi_{r_1 \dots r_{\lambda-1} i} \frac{\partial X_{r_\lambda}}{\partial x_i} \right) = 0$$

identisch erfüllen. Auf kanonische Systeme angewendet, ergibt sich ein Theorem, das als Spezialfall einen Satz von Lagrange in sich schließt.

Den Funktionen $M [q_{\alpha_1} q_{\alpha_2} \dots q_{\alpha_\nu} p_{\beta_1} p_{\beta_2} \dots p_{\beta_\mu}]$ von $t, q_1 \dots q_m, p_1 \dots p_m$ stehen jetzt die Funktionen

$$\xi [q_{\alpha_1} q_{\alpha_2} \dots q_{\alpha_\nu} p_{\beta_1} p_{\beta_2} \dots p_{\beta_\mu}] \quad (\nu + \mu = 2k)$$

von $t, q_1 \dots q_m, p_1 \dots p_m$ gegenüber und den Funktionen $M_{\alpha_1 \alpha_2 \dots \alpha_k}$ die Funktionen $\xi_{\alpha_1 \alpha_2 \dots \alpha_k}$.

Man erhält dann durch eine analoge Überlegung wie früher, daß

$$\frac{d \sum \xi_{\alpha_1 \alpha_2 \dots \alpha_k}}{dt} = 0$$

ist, d. h.

$\sum \xi_{\alpha_1 \alpha_2 \dots \alpha_k} = \text{konst.}$ ist ein Integral des kanonischen Systems (3).

Kennt man alle $2m$ Integrale $\varphi_i = c_i \quad i = 1, 2, \dots, 2m$ von (3), so kann man eine Variationslösung zweiter Ordnung ξ_α sofort bilden. Man braucht nur $q_1 \dots q_m, p_1 \dots p_m$ als Funktionen der Konstanten $c_1 \dots c_{2m}$ und t auszudrücken, dann ist, wenn c und c' zwei dieser Konstanten bedeuten,

$$\frac{\partial (q_\alpha p_\alpha)}{\partial (c \ c')}, \quad \alpha = 1, 2, \dots, k \text{ eine Variationslösung}$$

und $\sum_\alpha \frac{\partial (q_\alpha p_\alpha)}{\partial (c \ c')} = \text{konst.}$ ist Integral des kanonischen Systems (3) (Theorem von Lagrange).

Was die Literatur zu diesen und ähnlichen Betrachtungen anlangt, sei auf „Th. de Donder, Théorie des invariants intégraux“ Paris 1928, verwiesen, wo mit Hilfe des „Integralkalküls“ eine Vereinfachung des obenstehenden Beweises durchgeführt ist.

(Eingegangen: 21. III. 1928).

Affine Differentialgeometrie der Strahlensysteme I.

Von W. Haack in Stuttgart.

In einer langen Kette von Arbeiten haben vor allem W. Blaschke, G. Pick, L. Berwald, A. Winternitz u. a. die affine Differentialgeometrie der Kurven und Flächen aufgebaut. Im zweiten Bande von Blaschkes Vorlesungen über Differentialgeometrie¹⁾ hat sie eine einheitliche Zusammenfassung erfahren. Im folgenden soll der Aufbau einer affinen Differentialgeometrie der Strahlensysteme versucht werden. Als Grundlage der analytischen Darstellung werden Plückersche Linienkoordinaten verwendet. Doch ist von vornherein klar, daß sich eine Affingeometrie der Strahlensysteme nicht ausschließlich im Plückerschen Geradenraum bewegen kann, sondern sich auch mit Gebilden des Punkt- und Ebenenraumes beschäftigen wird. Es sei z. B. auf die Mittenfläche des Strahlensystems verwiesen. Man wird daher einen analytischen Apparat fordern, der solche invarianten Flächen direkt punktweise darzustellen gestattet, ohne auf den Punkt als Träger eines Strahlenbüschels oder -bündels angewiesen zu sein. Gleichzeitig bleibt wünschenswert, daß man in der Darstellung der Flächen den Anschluß an die affine Flächen-theorie gewinnt.

Die Affingeometrie macht eine Untergruppe der projektiven Geometrie aus, bei der ein kontravarianter Vektor, die uneigentliche Ebene, invariant bleibt. Hier erblickt man die erste Schwierigkeit, die einer Affingeometrie im Plückerschen Geradenraum entgegen-tritt. Die invariante uneigentliche Ebene läßt sich nur durch ihre Geraden darstellen und man müßte zur Abgrenzung der Affingeometrie eine ausgezeichnete, ausgeartete lineare Kongruenz invariant halten. Statt dessen ist es ratsam, aus dem Plückerschen Geradenraum herauszutreten und die uneigentliche Ebene W (so wird sie im folgenden stets benannt) als kontravarianten Vektor mit seinen vier homogenen Koordinaten einzuführen, also gewissermaßen den Plückerschen Geradenraum um diesen ausgezeichneten kontravarianten Vierervektor zu erweitern. Für die so erweiterten Plückerschen Koordinaten lassen sich einfache Rechenregeln angeben, die

¹⁾ Berlin, Verlag von Jul. Springer, 1923.

im wesentlichen auf invariantentheoretische Arbeiten von Mertens und Weitzenböck zurückgehen und unter dem Namen Mertens' Ketten bekannt sind²⁾.

Die vorliegende erste Mitteilung gibt zunächst den formalen Apparat, indem ein begleitendes Grundtetraeder konstruiert und die zugehörigen Ableitungsgleichungen und Integrierbarkeitsbedingungen aufgestellt werden. Man erkennt aus diesen, daß das Strahlensystem vollständig bis auf raumtreue Affinitäten bestimmt wird von einer quadratischen und einer kubischen Grundform.

Als erste geometrische Anwendung des formalen Apparates werden die Regelflächen des Strahlensystems untersucht, d. h. diejenigen geradlinigen Flächen, deren Erzeugende Systemstrahlen sind. Es ergeben sich verschiedene Klassen von Regelflächen im Strahlensysteme, von denen die Nullflächen der kubischen Grundform — Symmetrieflächen genannt — und die geodätischen Regelflächen besonders hervorgehoben werden mögen. Der letzte Paragraph ist den Brennflächen und deren Ausartungen in Kurven und Torsen gewidmet. Das Kapitel bildet den Übergang zur affinen Krümmungstheorie der Strahlensysteme, die sich mit der Diskussion der simultanen Invarianten der beiden Grundformen befaßt. Dies bleibt jedoch einer folgenden Mitteilung vorbehalten, die zugleich noch einiges über Ribeaucoursche Strahlensysteme bringen wird.

An dieser Stelle sei Herrn W. Blaschke und Herrn G. Thomson für Anregung und Ratschläge herzlicher Dank ausgesprochen. Besonderen Dank schulde ich Herrn B. L. van der Waerden, der mich auf Mertens' Ketten aufmerksam machte.

§ 1.

Die Plückerschen Koordinaten einer Geraden erhält man bekanntlich aus den Matrizien:

$$(1) \quad \left\| \begin{array}{cccc} x_1 & x_2 & x_3 & x_4 \\ y_1 & y_2 & y_3 & y_4 \end{array} \right\| \quad \text{oder} \quad \left\| \begin{array}{cccc} u^1 & u^2 & u^3 & u^4 \\ v^1 & v^2 & v^3 & v^4 \end{array} \right\|;$$

dabei sind x_i ; y_i die homogenen Koordinaten zweier Punkte; u^i , v^i diejenigen zweier Ebenen. Im ersten Fall bezeichnen wir die Koordinaten der Geraden mit p_{ik} , im zweiten mit p^{ik} . Dann sind also:

$$(2) \quad p_{ik} = \sigma (x_i y_k - y_i x_k) \quad \text{und} \quad p^{ik} = \rho (u^i v^k - v^i u^k).$$

Diese verschiedenen Darstellungen ein und derselben Geraden p_{ik} und p^{ik} unterscheiden sich nur in der Reihenfolge der Koordinaten, da ja ein allen Komponenten gemeinsamer Faktor vernachlässigt werden darf. Deshalb kann im folgenden eine Gerade einfach

²⁾ Vgl. Weitzenböck, Komplex-Symbolik, Leipzig 1908 (Sammlung Schubert, Bd. 57, S. 27 ff.).

durch einen Sechservektor p (d. h. einen kleinen deutschen Buchstaben) dargestellt werden, dessen Komponenten entweder als p_{ik} oder als p^{ik} zu denken sind.

Ein Sechservektor p mit sechs willkürlichen Komponenten kann geometrisch stets als linearer Geradenkomplex gedeutet werden. Als Gerade dagegen nur dann, wenn dieser Komplex ein spezieller ist, d. h. wenn die Komponenten von p der Plückerschen Identität genügen:

$$(3) \quad \frac{1}{4} p^2 = p_{12} p_{34} + p_{13} p_{42} + p_{14} p_{23} = 0.$$

Mit Hilfe dieser Identität wird das Skalarprodukt zweier Sechservektoren definiert als Polarenbildung der quadratischen Gleichung (3), also das Skalarprodukt der beiden Vektoren p und q :

$$(4) \quad \frac{1}{2} p q = p_{12} q_{34} + p_{13} q_{42} + p_{14} q_{23} + p_{34} q_{12} + p_{42} q_{13} + p_{23} q_{14}.$$

Ist $p^2 = 0$ und $q^2 = 0$, so besagt $p q = 0$, daß sich die Geraden p und q schneiden. Ist dagegen $p^2 = 0$ und $q^2 \neq 0$, dann sagt $p q = 0$ aus, daß die Gerade p im Komplex q liegt.

Jetzt zur Definition des Strahlensystems:

Ein Stück M einer zwei-dimensionalen Mannigfaltigkeit des vier-dimensionalen Plückerschen Geradenraumes heie ein Strahlensystem, wenn die 6 homogenen Koordinaten durch die Gleichungen

$$p_{ik} = f_{ik}(u, v)$$

gegeben sind, wo die f_{ik} im Gebiete M des Parameterraumes (u, v) n -mal stetig differenzierbare Funktionen sind, die in M stets der Plückerschen Identität genügen und deren Funktionalmatrix

$$\left\| \begin{array}{cc} f_{ik} & \frac{\partial f_{ik}}{\partial u} \quad \frac{\partial f_{ik}}{\partial v} \end{array} \right\|$$

stets den Rang drei hat.

Da die Komponenten p_{ik} homogene Koordinaten sind, ändert sich das Strahlensystem nicht, wenn man statt $p_{(u, v)}$

$$\hat{p}_{(u, v)} = \lambda_{(u, v)} p_{(u, v)} \quad (\lambda = \text{Skalar})$$

setzt; vorausgesetzt, daß auch $\lambda_{(u, v)}$ genügend oft differenzierbar ist. Man wird also durch geeignete Wahl der Funktion $\lambda_{(u, v)}$ eine Normierung der homogenen Koordinaten einführen. Die Bestimmung dieser Normierung lassen wir zunächst offen und nehmen an, $p_{(u, v)}$ sei ein irgendwie normiertes Strahlensystem.

Jetzt ist ein geeigneter Formelapparat für die affine Differentialgeometrie der Strahlensysteme aufzustellen; zunächst zu den Ketteninvarianten, die uns in naturgemäßer Weise zu ko- und kontravarianten Vierervektoren führen.

§ 2.

Ist ein Strahlensystem durch den Plückerschen Vektor $p_{(u,v)}$ nach der Definition des vorigen Paragraphen gegeben und soll die Umgebung eines Strahles $p_{(u_0, v_0)}$ betrachtet werden, so wird man zunächst nach invarianten Punkten fragen, die auf $p_{(u_0, v_0)}$ selbst liegen. Bekanntlich sind dies die Brennpunkte, der uneigentliche Punkt und der Mittelpunkt zwischen den Brennpunkten; oder allgemeiner: jeder Punkt des Strahles, wenn er durch sein Doppelverhältnis mit dem uneigentlichen Punkt und den beiden Brennpunkten gegeben ist. Man sieht weiter, wie man zu beliebig vielen, affin-invariant mit dem Strahlensystem verbundenen Flächen gelangt, wie z. B. den Brennpunkten, der Mittenfläche oder allgemeiner allen Flächen, die die Brennpunktsstrecke in konstantem Verhältnis teilen. In einer affinen Differentialgeometrie der Strahlensysteme wird man nicht umhin können, solche Flächen zu untersuchen. Aber schon die analytische Darstellung der Flächen bereitet Schwierigkeiten. Es ist zu zeigen, wie man, ausgehend von Plückers Linienkoordinaten $p_{(u,v)}$, zu einer solchen Darstellung der Flächen gelangt, die es gestattet, unmittelbar Blaschkes Formeln der affinen Flächentheorie anzuwenden.

Im folgenden sind $U, V, W \dots$ bzw. $u^i; v^i; w^i \dots$ kontravariante und $x; y \dots$ bzw. $x_i; y_i \dots$ kovariante Vierer-Vektoren; also i läuft von 1 bis 4. Hat man eine Ebene U und eine Gerade p (d. h. p_{ik} oder p^{ik}), so bildet man:

$$(1) \quad u^i p_{ik} = x_k \quad [\text{über } i \text{ summiert}^3]$$

d. h. $u^i p_{ik}$ gibt einen kovarianten Vektor; dieser stellt, wenn p eine Gerade ist, den Schnittpunkt von p mit der Ebene U , wenn p ein linearer Komplex ist, den Pol der Ebene U in bezug auf den Komplex p dar.

Ganz analog hat man:

$$(2) \quad y_i p^{ik} = v^k;$$

also $y_i p^{ik}$ ist ein kontravarianter Vektor, der für $p^2 = 0$ die durch den Punkt y und die Gerade p gehende Ebene angibt, dagegen für $p^2 \neq 0$ die Polarebene des Punktes y in bezug auf den linearen Komplex p .

Aus (1) und (2) sieht man, wie die Ketten

$$(3) \quad u^i p_{ik} q^{kl} = v^l \quad \text{und} \quad y_i p^{ik} q_{kl} = z_l$$

³⁾ Über Indizes, die in einem Ausdruck sowohl oben wie unten auftreten, wird wie üblich stets summiert.

geometrisch zu deuten sind. Ferner ist evident, daß man, ohne Mißverständnisse fürchten zu müssen, die Indizes fortlassen kann. Z. B. schreibt man für die letzten beiden Ketten einfach:

$$(4) \quad Upq = V \quad \text{und} \quad ypq = z,$$

Ausdrücke der Form

$$Upqx = u^i p_{ik} q^{ki} x_i$$

sind Invarianten. Ihr Verschwinden sagt aus: Der Punkt x liegt in der Ebene Upq . Ebenso übersieht man den Sinn anderer Ketten-Invarianten wie

$$UpV \quad \text{oder} \quad UpqrV \quad \text{usw.}$$

Wir geben einige Rechenregeln für die Ketten, die sich aus dem Vorhergehenden mühelos ableiten lassen⁴⁾:

1. $Upqr \dots = \pm \dots rqp U$, je nachdem die Anzahl der Sechservektoren der Kette gerade oder ungerade ist; denn es ist $p^{ik} = -p^{ki}$. Das Entsprechende gilt für die Ketten: $xpqr \dots = \pm \dots rqp x$.

2. $Upqr \dots = -Upqr \dots - \frac{1}{2} (pq) Ur \dots$ und analog: $xpqr \dots = -xpqr \dots - \frac{1}{2} (pq) xr \dots$, dabei ist (pq) das in § 1 definierte Skalarprodukt zweier Sechservektoren. Diese Regel gibt einen Ersatz für das kommutative Gesetz.

Die Ketten zeigen, wie man, ausgehend von Plückerschen Linienkoordinaten, zu homogenen Punkt- und Ebenenkoordinaten gelangen kann; jedoch ist dazu notwendig, daß man zunächst schon mindestens einen ko- oder kontravarianten Vektor kennt, daß also mindestens die Koordinaten eines Punktes oder einer Ebene bekannt sind. Man kann sagen: Erweitert man den Plückerschen Geradenraum um einen einzigen Punkt oder eine Ebene, so kann man sofort mit Hilfe der Ketten sämtliche Punkte und Ebenen in ihren gewöhnlichen homogenen Koordinaten darstellen. In der Affin-Geometrie ist nun die uneigentliche Ebene, die stets mit W bezeichnet werden soll, ausgezeichnet. Es liegt also auf der Hand, den Plückerschen Geradenraum um diesen Vierervektor zu erweitern. Dann lassen sich die homogenen Koordinaten aller Punkte und Ebenen durch Ketten, die lediglich mit W und einer Anzahl Sechservektoren gebildet sind, leicht angeben. Zum Beispiel wird Wp der uneigentliche Punkt der Geraden p . Im folgenden treten in der Hauptsache Invarianten der Form

$$WpqrW$$

auf, deren Verschwinden bedeutet, daß der Punkt $Wpqr$ in der uneigentlichen Ebene liegt, oder daß die Ebene Wpq den Punkt

⁴⁾ Vgl. über die Ketten etwa: Weitzenböck, Komplex-Symbolik, Leipzig 1908, pag. 27 ff.

r W enthält, oder auch daß der Punkt Wp in der Ebene qr W liegt, was auch heißt: daß die Gerade p der Ebene qr W parallel ist.

§ 3.

Jedes Strahlensystem $p_{(u,v)}$ enthält im allgemeinen zwei Scharen von Torsen (= abwickelbaren Flächen), deren Differentialgleichung wir als quadratische Grundform wählen:

$$(1) \quad \Phi = p_u^2 du^2 + 2 p_u p_v du dv + p_v^2 dv^2,$$

dabei sind p_u und p_v die partiellen Ableitungen des Vektors $p_{(u,v)}$, deren Komponenten also die partiellen Ableitungen der Komponenten von p sind. Für das Folgende werde zunächst angenommen, daß u und v Torsenparameter sind, daß also:

$$(2) \quad p_u^2 = 0 \text{ und } p_v^2 = 0 \text{ dagegen } p_u \cdot p_v = F \neq 0$$

sind. Dann sind p_u und p_v zwei Geraden. Wegen $p^2 = 0$ und daher auch $p p_u = 0$ und $p p_v = 0$ schneiden diese den Systemstrahl p . Die Schnittpunkte sind die Brennpunkte, ihr geometrischer Ort die Brennflächen des Strahlensystems. Die Tangentialebenen der Brennflächen heißen Brennebenen. Das Zusammenfallen der beiden Brennpunkte oder Brennebenen wird stets ausgeschlossen, weil dann die Grundform Φ ausartet. Ebenso wird im Folgenden vorausgesetzt, daß stets beide Brennpunkte im Endlichen liegen, daß also das Strahlensystem kein zylindrisches ist⁵⁾. Der Ort der Mittelpunkte der Brennpunktsstrecke ist die Mittenfläche des Strahlensystems.

Als Eckpunkte eines mit dem Strahlensystem affin-invariant verbundenen Grundtetraeders (= begleitendes Sechseck) wählen wir den Mittelpunkt (= Schnittpunkt von p mit der Mittenfläche) y und die uneigentlichen Punkte der Geraden p ; p_u und p_v . Ist wieder W der kontravariante Vierervektor der uneigentlichen Ebene, z. B. mit den Komponenten $\{0, 0, 0, 1\}$, dann werden die uneigentlichen Punkte von p , p_u und p_v nach § 2 gegeben durch

$$(3) \quad x = Wp, \quad x_u = Wp_u, \quad x_v = Wp_v.$$

Die Brennpunkte p und $\bar{p}$ sind die Schnittpunkte des Strahles p_u (bezw. p_v) mit der Ebene $Wp p_v$ (bezw. $Wp p_u$); d. h.

$$(4) \quad p = - Wp p_v p_u \text{ und } \bar{p} = + Wp p_u p_v.$$

Den Mittelpunkt zwischen $\bar{p}$ und p erhält man als vierten harmonischen Punkt zu x ; $\bar{p}$ und p . Es wird:

$$p - \bar{p} = - Wp p_v p_u - Wp p_u p_v$$

⁵⁾ Die hier ausgeschlossenen parabolischen und zylindrischen Strahlensysteme werden in einer späteren Arbeit untersucht werden.

$$= -W p p_v p_u + W p p_v p_u + \frac{1}{2} F W p = \frac{1}{2} F W p,$$

nach § 2. $\frac{1}{2} F \cdot W p$ ist aber der uneigentliche Punkt x . Man erhält also den Mittelpunkt y zu:

$$(5) \quad y = p + \bar{p} = W p_u p_v p - W p_v p_u p.$$

Vollkommen festgelegt sind nur die Punkte y und x des Grundtetraeders. Die Punkte x_u und x_v hängen nicht nur von den Parametern u und v ab, sondern sie ändern ihre Lage auch mit der Normierung der homogenen Plückerschen Linienkoordinaten des Systemstrahls p . Eine geometrische Deutung der Punkte x_u und x_v ist also erst möglich, wenn eine invariante Normierung gefunden wird.

Die den Ecken x , x_u , x_v , y des Grundtetraeders gegenüberliegenden Seitenflächen bezeichnen wir mit den kontravarianten Vektoren X , X^u , X^v , Y . Sie werden gegeben durch die Ketten

$$\begin{aligned} X &= -p_v p_u W + p_u p_v W \\ X^u &= p_v p W \\ X^v &= p p_u W \\ Y &= W \end{aligned} \quad (6)$$

Die Determinante des Grundtetraeders $(x x_u x_v y) = E_{12}$ schreibt sich:

$$(7) \quad \frac{1}{2} x X = x_u X^u = x_v X^v = \frac{1}{2} y Y = E_{12}^6).$$

Die homogenen Koordinaten eines beliebigen Punktes r werden in bezug auf das Grundtetraeder dargestellt durch den Ausdruck:

$$(8) \quad r = \alpha x + \beta x_u + \gamma x_v + \delta y,$$

dabei bestimmen sich seine „Koordinaten“ α , β , γ , δ durch Überschiebung der Gleichung (8) mit x , x_u , x_v , y aus der Proportion:

$$\alpha : \beta : \gamma : \delta = \frac{1}{2} r X : r X^u : r X^v : \frac{1}{2} r Y,$$

oder, wenn man sie einzeln bestimmt:

$$(9) \quad \alpha = \frac{r X}{2 E_{12}}, \quad \beta = \frac{r X^u}{E_{12}}, \quad \gamma = \frac{r X^v}{E_{12}}, \quad \delta = \frac{r Y}{2 E_{12}}.$$

⁶⁾ Hier sind u und v keine Summationsindizes!

Man sieht jetzt unmittelbar, daß der Ausdruck

$$(2) \quad \frac{W p p_u p_v W}{\sqrt{(p_u p_v)^2 - p_u^2 p_v^2}}$$

invariant ist gegenüber Parametertransformationen.

Wie ändert sich der Ausdruck (2), wenn man die sechs homogenen Koordinaten des Strahles $p_{(u, v)}$ mit einem Skalar $\lambda_{(u, v)}$ multipliziert?

Also man setzt:

$$\hat{p}_{(u, v)} = \lambda_{(u, v)} p_{(u, v)}.$$

Dann werden

$$\hat{p}_u = \lambda_u p + \lambda p_u$$

$$p_v = \lambda_v p + \lambda p_v$$

und daraus folgt:

$$\frac{W \hat{p} \hat{p}_u \hat{p}_v W}{\sqrt{(\hat{p}_u \hat{p}_v)^2 - \hat{p}_u^2 \hat{p}_v^2}} = \frac{\lambda W p p_u p_v W}{\sqrt{(p_u p_v)^2 - p_u^2 p_v^2}}.$$

Der Ausdruck (2) ist demnach nicht invariant gegenüber einer Änderung der homogenen Normierung. Durch die Forderung

$$(3) \quad \frac{W \hat{p} \hat{p}_u \hat{p}_v W}{\sqrt{(\hat{p}_u \hat{p}_v)^2 - \hat{p}_u^2 \hat{p}_v^2}} = 1$$

wird also eine Normierung der homogenen Plückersehen Koordinaten $p_{(u, v)}$ festgelegt, die gegen Parametertransformationen invariant ist und bei passender Wahl des Vorzeichens der Quadratwurzel eindeutig ist.

Im § 3 diene die Differentialgleichung der Torsen des Strahlensystems als quadratische Grundform:

$$\Phi = \hat{p}_u^2 du^2 + 2 \hat{p}_u \hat{p}_v du dv + \hat{p}_v^2 dv^2,$$

sie liefert uns jetzt den Grundtensor für den Ricci-Kalkül, indem wir setzen:

$$(4) \quad G_{ik} = \hat{p}_i \hat{p}_k \quad \left[\hat{p}_i = \frac{\partial \hat{p}}{\partial u^i}, \quad u^1 = u, \quad u^2 = v \right]$$

Der Tensorcharakter der G_{ik} ist bereits durch die Gleichungen (1) bewiesen. Jetzt läßt sich die Normierung (3) in einer anderen Form schreiben, wenn man den Tensor:

$$(5) \quad E_{ik} = W \hat{p}_i \hat{p}_k W$$

eingführt; nach § 2 ist offenbar: $E_{11} = E_{22} = 0$ und $E_{12} = -E_{21} \neq 0$.

Geschieht das Herauf- und Herunterziehen der Indizes in der üblichen Weise durch

$$E^{ik} = G^{il} G^{km} E_{lm},$$

so nimmt der Ausdruck (3) die Form an:

(6)

$$\frac{1}{2} E_{ik} E^{ik} = -1$$

Aus dieser Gleichung folgt, daß die kovarianten Ableitungen des Tensors E_{ik} bezüglich des Grundtensors G_{ik} verschwinden:

$$(7) \quad E_{ik \cdot l} = 0 \quad E^{ik \cdot l} = 0.$$

Das erkennt man durch folgende kleine Rechnung:

$$\frac{1}{2} E_{ik} G^{ir} G^{km} E_{rm} = -1$$

also:

$$E_{ik \cdot l} G^{ir} G^{km} E_{rm} + E_{ik} G^{ir} G^{km} E_{rm \cdot l} = 0.$$

Daraus folgt durch Zerlegung in Komponenten:

$$E_{12 \cdot l} \cdot E_{12} (G^{11} G^{22} - G^{12} G^{12}) = 0$$

d. h.

$$E_{12 \cdot l} = 0.$$

§ 5.

Im Anschluß an den Paragraphen 3 wenden wir uns jetzt wieder dem Grundtetraeder zu, um seine Koordinaten für allgemeine Parameter zu bestimmen. Da jetzt nicht mehr $p_i^2 = 0$ angenommen werden kann, müssen wir p_i als linearen Komplex deuten. Die drei Punkte des Tetraeders, die in der uneigentlichen Ebene liegen, sind nach § 3

$$Wp = x \text{ und } Wp_i = x_i,$$

dabei ist wieder Wp der uneigentliche Punkt der Geraden p , dagegen sind Wp_i die Pole der uneigentlichen Ebene in bezug auf die Komplexe p_i . Der vierte Eckpunkt des Grundtetraeders war der Mittelpunkt y ; die Bestimmung seiner Koordinaten gestaltet sich jetzt in allgemeinen Parametern komplizierter. Zunächst werden die Brennpunkte p und $\bar{p}$ des Systemstrahles p ermittelt, dann ist der Mittelpunkt y bekanntlich, von einem unwesentlichen Faktor abgesehen, nach § 3:

$$(1) \quad y = \frac{1}{2} (p + \bar{p}).$$

Es seien p_{1_0} und p_{2_0} die Ableitungen nach den Torsenparametern, dann sind die Brennpunkte nach § 3:

$$p = Wp p_{1_0} p_{2_0} \text{ und } \bar{p} = -Wp p_{2_0} p_{1_0}.$$

Die Differentialgleichung der Torsen ist $G_{ik} du^i du^k = 0$; demnach werden:

$$\begin{aligned} p_{1_0} &= p_1 G_{22} + p_2 (-G_{12} - D) \\ p_{2_0} &= p_1 G_{22} + p_2 (-G_{12} + D) \end{aligned} \quad (D = \sqrt{G_{12}^2 - G_{11} G_{22}}).$$

Setzt man diese Werte in die Ausdrücke der Brennpunkte ein, so wird:

$$\left. \begin{aligned} p \\ \bar{p} \end{aligned} \right\} = \pm G_{22} Wp p_1 p_1 \mp G_{12} (Wp p_1 p_2 + Wp p_2 p_1) \pm G_{11} Wp p_2 p_2 + D (Wp p_1 p_2 - Wp p_2 p_1).$$

Nun ist aber nach § 2:

$$Wp p_1 p_1 = -Wp p_1 p_1 - \frac{1}{2} Wp \cdot G_{11} = -\frac{1}{4} G_{11} Wp$$

$$Wp p_2 p_2 = -\frac{1}{4} G_{22} Wp$$

$$Wp p_2 p_1 = -Wp p_1 p_2 - \frac{1}{2} G_{12} Wp.$$

Man kann also die Brennpunkte schreiben:

$$\left. \begin{aligned} p \\ \bar{p} \end{aligned} \right\} = \pm \frac{1}{2} (G_{12}^2 - G_{11} G_{22}) Wp + D (Wp p_1 p_2 - Wp p_2 p_1).$$

Infolge der Normierung [Gl. (6) § 4] $\frac{1}{2} E_{ik} E^{ik} = -1$ ist:

$$\frac{1}{D} \cdot E^{12} (G_{11} G_{22} - G_{12}^2) = 1,$$

da aber $D^2 = G_{12}^2 - G_{11} G_{22}$ ist, so wird

$$E^{12} \cdot D = -1$$

und wir können schreiben:

$$\left\{ \begin{matrix} \bar{p} \\ p \end{matrix} \right\} = \pm \frac{1}{2} D^2 W p - E^{12} D^2 (W p p_1 p_2 - W p p_2 p_1).$$

Oder, wenn wir den Faktor D^2 vernachlässigen:

$$(2) \quad \left\{ \begin{matrix} \bar{p} \\ p \end{matrix} \right\} = \pm \frac{1}{2} W p - E^{ik} W p p_i p_k.$$

Für den Mittenpunkt y folgt daraus nach Gl. (1)

$$(3) \quad \boxed{y = -\frac{1}{2} E^{ik} W p p_i p_k.}$$

Setzt man noch $W p = x$, dann werden die Brennpunkte

$$(4) \quad \boxed{\left\{ \begin{matrix} \bar{p} \\ p \end{matrix} \right\} = y \pm \frac{1}{4} x.}$$

Ein beliebiger Punkt r wird, bezogen auf das Grundtetraeder, dargestellt durch:

$$(5) \quad r = a x + b^i x_i + c y.$$

Seine Koordinaten a, b^i, c bestimmen sich analog wie im § 3 unter Heranziehung der den Ecken x, x_i, y gegenüberliegenden Seitenflächen X, X^i, Y des Grundtetraeders. Offenbar ist

$$Y = W.$$

Die Determinante des Grundtetraeders wird jetzt

$$(6) \quad y Y = \frac{1}{2} x_i X^i = x X = -\frac{1}{2} E_{ik} E^{ik} = 1.$$

Daraus folgen für X und X^i die Ausdrücke:

$$(7) \quad \begin{aligned} X &= -p_i p_k W \cdot \frac{1}{2} E^{ik} \\ X^i &= -p_r p W E^{ir}. \end{aligned}$$

Schließlich werden die „Koordinaten“ des Punktes r :

$$(5a) \quad a = rX, \quad b^i = rX^i, \quad c = rY.$$

Dazu findet man unmittelbar die Tafel der Skalarprodukte:

$$(8) \quad \begin{cases} xX = 1 & x_i X = 0 & yX = 0 \\ xX^i = 0 & x_i X^i = 2 & yX^i = 0 \\ xY = 0 & x_i Y = 0 & yY = 1 \end{cases}$$

Zum Schluß dieses Paragraphen mögen noch die Ausdrücke einer beliebigen, die Strahlen des Systems schneidenden Fläche angeführt werden. Eine solche Fläche wird den Strahl $p_{(u,v)}$ in einem Punkte $f_{(u,v)}$ schneiden. Bezogen auf das Grundtetraeder muß sich dann $f_{(u,v)}$ schreiben:

$$f_{(u,v)} = \beta y + \alpha x.$$

Um $f_{(u,v)}$ jedoch invariant darzustellen, ist die Funktion $\frac{\alpha_{(u,v)}}{\beta_{(u,v)}}$

zu normieren; dies erreicht man dadurch, daß man $\frac{\alpha}{\beta} = 1$ setzt, wenn f mit dem Brennpunkt p zusammenfällt. Dann schreibt sich $f_{(u,v)}$:

$$(9) \quad f_{(u,v)} = y + \frac{1}{4} \sigma_{(u,v)} x,$$

$\sigma_{(u,v)}$ gibt dann das Teilverhältnis, in welchem f die Strecke $\overline{yx}$ teilt.

Der Ausdruck (9) erfüllt in der Tat die Forderungen, die im Anfang des § 2 aufgestellt wurden: Es ist nämlich nach der Tafel der Skalarprodukte (8) $yY = 1$ und $xY = 0$, so daß man bekommt

$$fY = fW = 1.$$

Nimmt man den Vektor W der uneigentlichen Ebene in der Form $\{0, 0, 0, 1\}$ an, so werden die Komponenten von f : $\{f_1, f_2, f_3, 1\}$. Die homogenen Koordinaten von f stimmen also infolge der Normierung $\frac{1}{2} E^{ik} E_{ik} = -1$ mit den inhomogenen Parallelkoordinaten überein, so daß man die Formeln der affinen Flächentheorie (Blaschke II) unmittelbar auf die Fläche $f_{(u,v)}$ anwenden kann.

§ 6.

Die Ableitungsgleichungen und Integrierbarkeitsbedingungen bilden das Grundgerüst jeder Differentialgeometrie; sie mögen jetzt für die affine der Strahlensysteme zusammengestellt

werden. Bedeuten, wie üblich, die Indizes die kovarianten Ableitungen bezüglich des Grundtensors G_{ik} , dann versteht man unter den Ableitungsgleichungen die Ausdrücke der Punkte

$$x_{kl} \text{ und } y_r$$

bezogen auf das Grundtetraeder x, x_i, y .

Zur Abkürzung bezeichnen wir die Tensoren:

$$(1) \quad \begin{aligned} A_{ik} &= \frac{1}{2} W p_{ik} p_i p_m W E^{im} \\ B_{ikl} &= -W p_{ik} p_l p. \end{aligned}$$

Infolge der Normierung $\frac{1}{2} E_{ik} E^{ik} = -1$ wird der Tensor B_{ikl}

symmetrisch, denn es ist offenbar $B_{ikl} = B_{kli}$; außerdem aber, wegen $E_{ik..i} = 0$ (§ 4) auch $B_{kii} = B_{ilk}$.

Den Punkt x_{kl} kann man nach § 5, Gl. (5), darstellen durch:

$$x_{kl} = a_{kl} x + b_{kl}^i x_i + c_{kl} y,$$

dabei sind nach § 5, Gl. (5a)

$$a_{kl} = x_{kl} X, \quad b_{kl}^i = x_{kl} X^i, \quad c_{kl} = x_{kl} Y.$$

Nach § 5 ist $x_{kl} = W p_{kl}$, deshalb wird nach § 5, Gl. (7)

$$a_{kl} = -W p_{kl} p_i p. W \cdot \frac{1}{2} E^{is} = -A_{kl}$$

und ebenso:

$$\begin{aligned} b_{kl}^i &= -W p_{kl} p_r p W E^{ir} = B_{klr}. E^{ir} \\ c_{kl} &= W p_{kl} W = 0. \end{aligned}$$

Und es kommt schließlich:

(2)

$$x_{kl} = -A_{kl} x + E^{ir} B_{klr} x_i.$$

Etwas schwieriger gestaltet sich die Bestimmung von y_r . Wir gehen in derselben Weise vor und setzen:

(3)

$$y_r = a_r x + b_r^i x_i + c_r y.$$

Zunächst wird nach § 5, Gl. (5a):

$$c_r = y_r W = 0,$$

da $y Y = y W = 1$ ist und die Ableitungen von W verschwinden. Nach § 5, Gl. (3), wird:

$$y_r = -\frac{1}{2} E^{ik} \{ W p_r p_i p_k + W p p_{ir} p_k + W p p_i p_{kr} \};$$

daraus kann man b_r^i und a_r nach § 5, Gl. (5a), finden. Jedoch kommt man schneller zum Ziel, wenn man setzt:

$$b_r^i = y_r X^i = -y X_r^i$$

und

$$a_r = y_r X = -y X_r;$$

denn dann wird nach § 5

$$b_r^i = -\frac{1}{2} E^{ik} E^{lm} \{ W p p_i p_k p_{mr} p W + W p p_i p_k p_m p_r W \}.$$

Die beiden Ketten sind nach den Regeln des § 2 auszuwerten. Man vertauscht z. B. in der ersten Kette das letzte p dreimal, so daß dann die beiden in der Kette auftretenden p nebeneinander stehen:

$$\begin{aligned} W p p_i p_k p_{mr} p W &= - W p p_i p_k p p_{mr} W - \frac{1}{2} (p p_{mr}) W p p_i p_k W \\ &= + W p p_i p p_k p_{mr} W - \frac{1}{2} (p p_{mr}) W p p_i p_k W \quad (\text{da } p p_k = 0) \\ &= - W p p p_i p_k p_{mr} W - \frac{1}{2} (p p_{mr}) W p p_i p_k W \\ &= - \frac{1}{2} (p p_{mr}) W p p_i p_k W = + \frac{1}{2} G_{mr} E_{ik}. \end{aligned}$$

Ähnlich behandelt man die zweite Kette; ohne die elementare Rechnung anzuführen, findet man

$$b_r^i = \frac{1}{2} E^{im} G_{mr} - \frac{1}{4} E^{lm} G_{mr} = + \frac{1}{4} E^{lm} G_{mr}.$$

Für a_r erhält man die Ketten:

$$\begin{aligned}
 a_r &= -\frac{1}{4} E^{ik} E^{lm} \{ W p_i p_k p_{lr} p_m W + W p_i p_k p_l p_{mr} W \} \\
 &= -\frac{1}{2} E^{ik} E^{lm} W p_i p_k p_l p_{mr} W \quad (\text{da } p_m p_{lr} = 0).
 \end{aligned}$$

Daraus folgt schließlich:

$$a_r = \frac{1}{4} E^{ik} E^{lm} G_{km} B_{ilr}.$$

Es ist aber $E^{ik} E^{lm} G_{km} = G^{il}$, so daß man für a_r auch kürzer schreiben kann:

$$a_r = \frac{1}{4} B_{ir}^i.$$

Man hat also erhalten:

$$(4) \quad \boxed{y_r = \frac{1}{4} B_{ir}^i x + \frac{1}{4} E^{ik} G_{kr} x_i.}$$

Die Integrierbarkeitsbedingungen ergeben sich jetzt unmittelbar aus den Gleichungen (2) und (4). Man hat nämlich (Blaschke, II, pag. 151/152)

$$x_{kls} - x_{ksl} = R_{km \cdot sl} G^{mi} x_i,$$

wo $R_{km \cdot sl}$ der Riemannsche Krümmungstensor der quadratischen Grundform $G_{ik} du^i du^k$ ist; und außerdem:

$$y_{rs} - y_{sr} = 0.$$

Man kommt schließlich zu den Integrierbarkeitsbedingungen:

$$(5) \quad \boxed{
 \begin{aligned}
 E^{rs} B_{ir \cdot s}^i - G^{ik} A_{ik} &= 0 \\
 E^{sl} A_{sl \cdot i} + E^{ir} E^{sj} B_{kr \cdot s} A_{il} &= 0 \\
 A_{kp} &= S G_{kp} - E^{\sigma q} B_{kp \sigma \cdot q} - E^{\sigma q} E^{ir} B_{ip \cdot q} B_{kr \sigma}
 \end{aligned}
 }$$

Hier ist S das Gaußsche Krümmungsmaß der quadratischen Grundform.

Aus den Integrierbarkeitsbedingungen (5) erkennt man, daß der Tensor A_{kp} durch die Tensoren G_{ik} und B_{ikl} völlig bestimmt ist,

während infolge der Normierung E_{ik} durch G_{ik} gegeben ist. Es folgt also der Satz:

Das Strahlensystem $p_{(u,v)}$ ist durch die quadratische Grundform $G_{ik} du^i du^k$ und die kubische Form $B_{ikl} du^i du^k du^l$ vollständig bis auf Affinitäten bestimmt.

Außer der kubischen Form $B_{ikl} du^i du^k du^l$ werden wir im nächsten Paragraphen eine andere kubische Form finden, die dann als kubische Grundform dienen soll.

§ 7.

Regelflächen im Strahlensystem.

Eine beliebige Fläche mit geradlinigen Erzeugenden (= Regelfläche) läßt sich darstellen in der Form:

$$(1) \quad r_{(\sigma t)} = a_{(t)} + \sigma b_{(t)},$$

wo durch $a_{(t)}$ und $b_{(t)}$ zwei Kurven gegeben sind und t und σ die Parameter der Fläche bedeuten. Berechnet man die Asymptotenlinien der Regelfläche, so wird man auf eine Riccatische Differentialgleichung geführt, die die Form hat:

$$\frac{d\sigma}{dt} = A\sigma^2 + B\sigma + C.$$

Wenn es gelingt, auf der Regelfläche den Parameter σ so zu wählen, daß ihm eine affin-invariante geometrische Bedeutung zukommt, dann wird durch jede der Gleichungen $A=0$, $B=0$ und $C=0$ eine affin-invariante Klasse von Regelflächen definiert. Dabei ist im besonderen die Klasse $B=0$ durch eine Symmetrie zur Kurve $\sigma=0$ (also $a_{(t)}$) ausgezeichnet.

Bestimmt man im Strahlensystem durch den Strahl, $p_{(u,v)}$ eine Fortschreitungsrichtung mit $\frac{du^i}{dt} = \dot{u}^i$, so werden die Punkte r einer Regelfläche im Strahlensystem gegeben durch

$$(2) \quad r_{(t\sigma)} = y_{(t)} + \frac{1}{4} \sigma x_{(t)}.$$

Die Parameterlinien $t = \text{konst.}$ sind die erzeugenden Systemstrahlen; während die Kurven $\sigma = \text{konst.}$ die Strecke zwischen dem Mittelpunkt y und dem Brennpunkt p in konstantem Verhältnis teilen. σ ist so invariant normiert, daß $\sigma = \pm 1$ je einen Brennpunkt gibt; d. h. σ ist ein im obigen Sinne invarianter Parameter.

Sind r_σ , r_t , $r_{\sigma\sigma}$, $r_{\sigma t}$, r_{tt} die gewöhnlichen partiellen Ableitungen des Vektors r nach den Parametern σ , t , dann folgt aus dem § 6, Gl. (2, 4)

$$(3) \quad \begin{cases} r_{\sigma} = \frac{1}{4} x \\ r_t = \frac{1}{4} B_{ii}^i \dot{u}^i x + \frac{1}{4} [E^{ik} G_{kr} \dot{u}^r + \sigma \dot{u}^i] x_i \end{cases}$$

und

$$(4) \quad \begin{cases} r_{\sigma\sigma} = 0 \\ r_{\sigma t} = \frac{1}{4} \dot{u}^i x_i \\ r_{tt} = \gamma x + (\alpha^i + \beta^i \sigma) x_i \end{cases}$$

dabei ist:

$$(4a) \quad \begin{cases} \alpha^s = \frac{1}{4} B_{ir}^i \dot{u}^r \dot{u}^s + \frac{1}{4} E^{ik} G_{kr} E^{sj} B_{ij} \dot{u}^r \dot{u}^s + \\ \quad \frac{1}{4} E^{sk} G_{kl} (\Gamma_{rt}^i \dot{u}^r \dot{u}^t + \ddot{u}^i) \\ \beta^s = \frac{1}{4} E^{sr} B_{ikr} \dot{u}^i \dot{u}^k + \frac{1}{4} (\Gamma_{ik}^s \dot{u}^i \dot{u}^k + \ddot{u}^s). \end{cases}$$

Die Größe γ wird im folgenden keine Verwendung finden und daher hier nicht explizit angeführt.

Die Differentialgleichung der Asymptotenlinien der Regelfläche r ist:

$$(5) \quad \mathfrak{G}_{11} d\sigma^2 + 2 \mathfrak{G}_{12} d\sigma dt + \mathfrak{G}_{22} dt^2 = 0,$$

wo die $\mathfrak{G}$ bedeuten:

$$\mathfrak{G}_{11} = |r, r_{\sigma} r_t r_{\sigma\sigma}|, \quad \mathfrak{G}_{12} = |r r_{\sigma} r_t r_{\sigma t}|, \quad \mathfrak{G}_{22} = |r r_{\sigma} r_t r_{tt}|.$$

Setzt man aus (3) und (4) die Werte in die Determinanten ein, so folgt:

$$\mathfrak{G}_{11} = 0$$

$$(6) \quad \mathfrak{G}_{12} = \frac{1}{4^2} G_{rs} \dot{u}^r \dot{u}^s$$

$$\mathfrak{G}_{22} = \frac{1}{4^2} [-G_{kr} \alpha^k \dot{u}^r + \sigma \{E_{ik} \alpha^k \dot{u}^i - G_{kr} \beta^k \dot{u}^r\} + \sigma^2 \cdot E_{ik} \beta^k \dot{u}^i].$$

Jeder der Koeffizienten von $\mathfrak{G}_{22}$, nämlich:

$$(7) \quad -G_{kr} \alpha^k \dot{u}^r \text{ und } E_{ik} \alpha^k \dot{u}^i - G_{kr} \beta^k \dot{u}^r \text{ und } E_{ik} \beta^k \dot{u}^i$$

definiert durch sein Verschwinden eine invariante Klasse von Regelflächen im Strahlensystem, deren Untersuchung einen Einblick in die geometrische Struktur des Strahlensystems gewähren wird.

Zunächst sei hier noch der Ausdruck der Affinnormalen der Regelfläche im Punkte r berechnet. Sie wird ihrer Richtung nach gegeben durch den uneigentlichen Punkt Δr , wo Δ den Beltrami'schen Differentiator bezüglich der Asymptotenlinien $\mathfrak{G}_{11}$, $\mathfrak{G}_{12}$, $\mathfrak{G}_{22}$ bedeutet^{a)}. Da $\mathfrak{G}_{11} = 0$, erhält man:

$$\Delta r = \frac{\partial \mathfrak{G}_{22}}{\partial \sigma} \cdot r_\sigma - 2 \mathfrak{G}_{12} r_{\sigma t};$$

durch Einsetzen aus Gl. (6) folgt dann:

$$(8) \quad \Delta r = 4 [(E_{ik} \alpha^k \dot{u}^i - G_{kr} \beta^k \dot{u}^r) + \sigma E_{ik} \beta^k \dot{u}^i] x - 2 G_{rs} \dot{u}^r \dot{u}^s \dot{u}^i x_i.$$

Den drei unter (7) aufgezählten Flächenklassen kommen folgende geometrischen Eigenschaften zu:

1) Es war

$$-G_{kr} \alpha^k \dot{u}^r = 0.$$

Daraus erhält man nach (4a)

$$(9) \quad -G_{sr} B_{it}^i \dot{u}^t \dot{u}^s \dot{u}^r - G_{st} E^{sj} G_{kr} E^{rk} B_{iq} \dot{u}^r \dot{u}^q \dot{u}^t - \\ - E_{rt} [\Gamma_{qt}^i \dot{u}^q \ddot{u}^t \dot{u}^r + \ddot{u}^t \dot{u}^r] = 0.$$

Aus den Gleichungen (5) und (6) ist ersichtlich, daß die Kurve $\sigma = 0$ auf diesen Regelflächen eine Asymptotenlinie ist. $\sigma = 0$ ist aber die Schnittkurve der Regelfläche mit der Mittenfläche des Strahlensystems. Die Regelflächen (9) werden also von der Mittenfläche in einer Asymptotenlinie geschnitten. Da nach Definition die Schmiegeebene einer Asymptotenlinie mit der Tangentialebene der Fläche zusammenfällt, sind die Regelflächen (9) auch durch den Satz geometrisch zu kennzeichnen: *Von den Regelflächen (9) werden auf der Mittenfläche diejenigen Kurven ausgeschnitten, deren Schmiegeebenen den Systemstrahl enthalten.*

2) Die Klasse

$$E_{ik} \beta^k \dot{u}^i = 0$$

nimmt nach (4a) die Gestalt an:

$$(10) \quad B_{jki} \dot{u}^j \dot{u}^k \dot{u}^i + E_{is} [\Gamma_{jk}^s \dot{u}^j \dot{u}^k \dot{u}^i + \ddot{u}^s \dot{u}^i] = 0.$$

^{a)} Vgl. Blaschke, II, pag. 105f.

Sie kennzeichnet diejenigen Regelflächen des Strahlensystems, die die uneigentliche Ebene in Geraden schneiden, d. h. die parabolischen Regelflächen; also gibt Gleichung (10) die Geraden der uneigentlichen Ebene. Dies folgt entweder wie in 1) aus den Gleichungen (5) und (6) oder direkt aus der Determinante

$$W p_i p_{tt} W = 0,$$

die unmittelbar auf die Gleichung (10) führt.

3) Am merkwürdigsten ist die dritte Flächenklasse:

$$(11) \quad E_{ik} \alpha^i \dot{u}^k - G_{kr} \beta^k \dot{u}^r = 0.$$

Führt man auch hier nach (5) und (6) die Werte α^i und β^k ein, so kommt:

$$(11a) \quad E^{ls} G_{si} B_{lkr} \dot{u}^i \dot{u}^k \dot{u}^r = 0.$$

Das ist eine Differentialgleichung erster Ordnung. Wir setzen zur Abkürzung

$$(12) \quad \boxed{E^{ls} G_{si} B_{lkr} = C_{ikr}}$$

und nennen die Form

$$(11b) \quad \boxed{\psi = C_{ikr} du^i du^k du^r}$$

die kubische Grundform des Strahlensystems und ihre Nullflächen (11a) die *Symmetrieflächen*. Da B_{ikl} symmetrisch ist, folgt:

$$C_{ikr} = C_{irk} \quad \text{und} \quad G^{ih} C_{ikr} = C^i_{.kr} = 0.$$

Jetzt zu den geometrischen Eigenschaften der Symmetrieflächen! Bekanntlich ist die quadratische Grundform der affinen Flächentheorie proportional zur Differentialgleichung der Asymptotenlinien einer Fläche. Die Symmetrieflächen haben die kennzeichnende Eigenschaft, daß ihre affine Grundform symmetrisch ist in bezug auf ihre Schnittkurve mit der Mittenfläche, d. h. in zwei Punkten einer Erzeugenden, die vom Mittenpunkt entgegengesetzt gleichen Abstand haben, sind die $\mathcal{G}_{11}, \mathcal{G}_{12}, \mathcal{G}_{22}$ entsprechend einander gleich. Offenbar ist diese Symmetrie unabhängig von der Normierung der Differentialgleichung der Asymptotenlinien, die notwendig ist, um zur affinen Grundform der Regelfläche zu gelangen⁹⁾. Aus der geometrischen Deutung der affinen Grundform läßt sich eine Deutung der Symmetrie konstruieren.

⁹⁾ Blaschke, II, pag. 102 und pag. 128.

In zwei benachbarten Punkten a, b einer Kurve $\sigma = \text{konst.} = k$ der Regelfläche werden die Tangenten an die Asymptotenlinien gezogen, die nicht die Erzeugenden der Regelfläche sind. Diese Tangenten schneiden die beiden Erzeugenden, auf denen a und b liegen, in zwei weiteren Punkten c, d . Der unendlich kleine Rauminhalt des Tetraeders $[a, b, c, d]$ sei V . Dann gilt als *Kennzeichen der Symmetriefflächen*: Das Volumen V längs der Kurve $\sigma = +k$ ist gleich demjenigen längs der Kurve $\sigma = -k$; d. h. es ist in Punkten, die symmetrisch zur Schnittkurve der Regelfläche mit der Mittenfläche auf derselben Erzeugenden liegen, das gleiche.

Eine andere, tiefer liegende Eigenschaft ergibt sich aus den *Affinnormalen der Symmetriefflächen*. Der uneigentliche Punkt Δr der Affinnormalen wird durch den Ausdruck (8) gegeben; für die Symmetriefflächen folgt daraus wegen (11)

$$(12) \quad \Delta r = 4 \sigma E_{ik} \beta^k u^i x - 2 (G_r, \dot{u}^r \dot{u}^s) \dot{u}^i x_i.$$

Für die Punkte $\sigma = 0$ erhält man:

$$(13) \quad [\Delta r]_{\sigma=0} = -2 (G_r, \dot{u}^r \dot{u}^s) \dot{u}^i x_i.$$

Dies sind drei Punkte der uneigentlichen Geraden $\overline{x_1 x_2}$, die man erhält, wenn für $\dot{u}^i$ die drei Lösungen der Gleichung (11) eingesetzt werden. Man kann daher den merkwürdigen Satz formulieren:

Die drei Affinnormalen, die im Mittenpunkte y auf den drei Symmetriefflächen durch p errichtet werden, liegen in einer Ebene. Sie heie die Normalenebene des Strahlensystems im Strahl p .

Die Ebene hat die Gleichungen:

$$n = y + \rho \dot{u}^i x_i.$$

Sie fällt offenbar mit einer Ebene des Grundtetraeders zusammen, so daß der Satz gilt: *Die Normalenebene fällt mit derjenigen Seitenfläche des Grundtetraeders zusammen, die dem uneigentlichen Punkt von p gegenüberliegt.* Dieser Satz gibt eine geometrische Interpretation der Normierung der homogenen Plückerschen Koordinaten

$$\frac{1}{2} E^{ik} E_{ik} = -1.$$

Angenommen, u und v seien Torsenparameter, so lassen sich drei Ebenen des Grundtetraeders unabhängig von der Normierung konstruieren; es sind die beiden Tangentialebenen an die Torsen durch p und die uneigentliche Ebene W . Von der vierten Tetraederebene ist zunächst nur der Mittenpunkt y fest, während die beiden Punkte Wp_i sich bei Änderung der Normierung auf den Geraden $\|Wp, Wp_i\|$ bewegen, so daß die zugehörige vierte Tetraederebene sich um den Punkt y dreht. Durch die Normierung $\frac{1}{2} E^{ik} E_{ik} = -1$

wird erreicht, daß diese noch unbestimmte Tetraederebene in die Normalenebene im Strahl p hineinfällt. Damit sind alle Teile des Grundtetraeders festgelegt.

Zwischen den soeben betrachteten drei Klassen von Regelflächen des Strahlensystems besteht noch ein geometrischer Zusammenhang, der zu einer anschaulichen *Konstruktion der Symmetrieflächen* führt.

Ausgehend von dem Systemstrahl $p_{(u,v)}$ gibt es nach jeder Richtung $\frac{dv}{du}$

1. eine Regelfläche, die von der Mitenfläche in einer Asymptotenlinie geschnitten wird, d. h. der Klasse 1, Gl. (9) angehört,

2. eine solche, die die uneigentliche Ebene längs einer Geraden schneidet, also der Klasse 2, Gl. (10), angehört.

In der willkürlich gewählten Richtung $\frac{dv}{du}$ haben die beiden Regelflächen zwei benachbarte Erzeugende p und $p + dp$ gemeinsam, während die nächste Nachbarerzeugende $p + dp + \frac{1}{2} d^2 p$ jeder Fläche im allgemeinen verschieden ist. Bildet man die Summe der Differentialgleichungen (9) und (10), so erhält man eine neue kubische Differentialform:

$$(14) \quad [B_{ikj} - G_{ik} B_{ij}^i - G_{sj} E^{st} G_{qt} E^{ta} B_{ikt}] \dot{u}^i \dot{u}^k \dot{u}^j = 0.$$

Diese Gleichung bestimmt gerade diejenigen drei Fortschreitungsrichtungen $\frac{dv}{du}$ im Strahlensystem durch einen Strahl $p_{(u,v)}$, für welche die Flächenklassen 1 und 2 drei benachbarte Erzeugende gemeinsam haben; oder in anderen Worten: für welche die Flächen der Klasse 1 in der Umgebung der Erzeugenden p parabolisch sind, d. h. ihre Schmiege- F_2 schneidet die uneigentliche Ebene in einer Geraden. Diese einfache Deutung der Nullflächen der Form (14) führt sofort zu den Symmetrieflächen. Wenn man zur Vereinfachung in (14) *Torsenparameter* einführt und die Normierung $\frac{1}{2} E^{ik} E_{ik} = -1$ beachtet, so läßt sich (14) schreiben:

$$(15) \quad C_{111} du^3 - C_{112} du^2 dv + C_{221} du dv^2 - C_{222} dv^3 = 0.$$

Dagegen werden die Symmetrieflächen in Torsenparametern durch die Gleichung gegeben:

$$C_{111} du^3 + C_{112} du^2 dv + C_{221} du dv^2 + C_{222} dv^3 = 0.$$

Die Nullrichtungen $\frac{dv}{du}$ dieser beiden Gleichungen unterscheiden sich lediglich um das Vorzeichen. Das bedeutet geometrisch, da

u, v Torsenparameter sind: Hat man, wie soeben angegeben, eine Fortschrittsrichtung der Gleichung (15) [bzw. (14)] konstruiert, und bestimmt man zu dieser und den beiden Torsenrichtungen die vierte harmonische Richtung durch den Strahl $p_{(u,v)}$, so fällt diese in eine Symmetrieffläche. Diese Konstruktion führt sehr anschaulich von den Geraden der uneigentlichen Ebene und denjenigen Kurven der Mittenfläche, deren Schmiegebene den Systemstrahl p enthält, zu den Symmetriefflächen. Die geometrischen Eigenschaften der Symmetriefflächen werden es rechtfertigen, gerade die Form $C_{ikl} du^i du^k du^l$ als kubische Grundform der affinen Differentialgeometrie der Strahlensysteme einzuführen.

§ 8.

**Geodätische Regelflächen, Bogenelement;
die Form $B_{ikl} du^i du^k du^l$.**

Im Strahlensystem seien wieder durch $\dot{u}^i = \frac{du^i}{dt}$ Regelflächen definiert. Man wird eine Regelfläche *geodätisch* nennen, wenn sie eine Extremale des Variationsproblems

$$(1) \quad \delta \int \sqrt{G_{ik} \dot{u}^i \dot{u}^k} dt = 0$$

darstellt. Das führt auf die bekannte Differentialgleichung II. Ordnung:

$$(2) \quad E_{ik} \dot{u}^i \ddot{u}^k + E_{ik} \Gamma_{rl}^k \dot{u}^i \dot{u}^r \dot{u}^l = 0.$$

Jetzt zurück zur Gleichung (1). Hier zeigt sich ein auffallender Zusammenhang mit den Gleichungen (5, 6) des vorigen Paragraphen. Jene Gleichungen gaben die zweite Gaußsche Grundform, also bis auf die Normierung die Blaschkesche I. Grundform einer durch das Strahlensystem gelegten Regelfläche. Nach W. Blaschke wird das Affin-Oberflächenelement einer Fläche gegeben durch

$$|\mathfrak{G}_{11} \mathfrak{G}_{22} - \mathfrak{G}_{12}^2|^{\frac{1}{4}} d\sigma dt,$$

wenn die $\mathfrak{G}$ die Koeffizienten der Differentialgleichung der Asymptotenlinien sind¹⁰⁾. Da für die Regelflächen nach dem vorigen Paragraphen $\mathfrak{G}_{11} = 0$ ist, findet man für die Affin-Oberfläche einer Regelfläche:

$$(3) \quad \Omega = \iint |\mathfrak{G}_{12}|^{\frac{1}{2}} d\sigma dt,$$

dabei ist nach § 7, Gleichung (6)

$$(4) \quad |\mathfrak{G}_{12}|^{\frac{1}{2}} = k |G_{ik} \dot{u}^i \dot{u}^k|^{\frac{1}{2}}, \quad (k = \text{Zahlenfaktor}),$$

¹⁰⁾ Über eine geometrische Deutung der Affinoberfläche Blaschke II, § 47.

wo $G_{ik} \dot{u}^i \dot{u}^k$ die erste Grundform des Strahlensystems ist. Man sieht: $\mathcal{G}_{12}$ ist eine Funktion von t allein, so daß das Doppelintegral (5) auch geschrieben werden kann:

$$(5) \quad \Omega = \int_{\sigma=\alpha}^{\beta} d\sigma \int_{t=a}^b \sqrt{G_{ik} \dot{u}^i \dot{u}^k} dt.$$

Für unsere Zwecke ist es vorteilhaft, die Integrationsgrenzen α und β von σ bestimmt festzulegen, und zwar wählen wir $\alpha = -1$ und $\beta = +1$. Dann wird:

$$(6) \quad \Omega = k \cdot \int_{t=t_1}^{t_2} \sqrt{G_{ik} \dot{u}^i \dot{u}^k} dt.$$

Dieses einfache Integral nennen wir ein *affines Regelstück des Strahlensystems*; es stellt den affinen Flächeninhalt eines Regelflächenstreifens dar von den Erzeugenden t_1 bis t_2 , der begrenzt wird von den beiden Berührungskurven der Regelfläche mit den beiden Brennflächen. Das *Regelement*

(7)

$$d\Omega = k \sqrt{G_{ik} \dot{u}^i \dot{u}^k} dt$$

gibt das vollständige Gegenstück zum Bogenelement der gewöhnlichen bewegungsinvarianten Flächentheorie, ebenso wie Gleichung (6) das Analogon zur Bogenlänge einer Flächenkurve ist.

Nachdem durch Gleichung (7) die erste Grundform (Torsen) des Strahlensystems geometrisch gedeutet ist als die Affinoberfläche des Regelflächenstreifens zwischen zwei benachbarten Systemstrahlen, fällt es nicht schwer, den *geodätischen Regelflächen* (1, 2) des Strahlensystems einen geometrischen Sinn zu geben, z. B. durch folgende Konstruktion: von dem Strahl p gehe man willkürlich zu einem Nachbarstrahl $p + p_i \dot{u}^i dt$, dann wähle man den nächsten Nachbarstrahl so, daß das affine Regelstück der entstehenden Regelfläche ein Extremum wird.

Bisher konnte keine geometrische Aussage über die kubische Form

(8)

$$B_{ikl} du^i du^k du^l$$

gemacht werden. Zunächst werden die Nullflächen dieser Form konstruiert, d. h. diejenigen Regelflächen des Strahlensystems, deren Differentialgleichung lautet

$$B_{ikl} \dot{u}^i \dot{u}^k \dot{u}^l = 0.$$

Gehen wir auf den vorigen Paragraphen zurück. Dort gab die Flächenklasse (2) diejenigen Regelflächen, die aus der uneigentlichen Ebene die Geraden ausschneiden. Sie genügten der Differentialgleichung (§ 7, Gl. 10)

$$(9) \quad E_{rs} [\dot{u}^r \dot{u}^s + \Gamma_{ik}^s \dot{u}^i \dot{u}^k \dot{u}^r] + B_{ikl} \dot{u}^i \dot{u}^k \dot{u}^l = 0.$$

Vergleicht man hiermit die Differentialgleichung der geodätischen Regelflächen (1), so sieht man, daß sich beide Differentialgleichungen nur um die Form $B_{ikl} \dot{u}^i \dot{u}^k \dot{u}^l$ unterscheiden. Daraus folgt unmittelbar eine Konstruktion der Regelflächen der Gleichung $B_{ikl} \dot{u}^i \dot{u}^k \dot{u}^l = 0$: Von dem Systemstrahl p gehe man in beliebiger Richtung zu einem Nachbarstrahl $p + d p$. Von diesem gehe man auf zwei Arten weiter zu einem zweiten Nachbarstrahl: 1. so, daß die entstehende Regelfläche geodätisch wird, und 2. so, daß die entstehende Regelfläche die uneigentliche Ebene in einer Geraden schneidet. Im allgemeinen wird man so zu zwei verschiedenen Strahlen kommen. Ändert man jetzt die Anfangsrichtung ($p + d p$), so findet man genau drei Richtungen von p aus derart, daß die beiden Konstruktionen des zweiten Nachbarstrahles $p + d p + \frac{1}{2} d^2 p$ zu demselben Strahl führen. Dies sind dann die Richtungen der Regelflächen $B_{ikl} \dot{u}^i \dot{u}^k \dot{u}^l = 0$. Man kann daher nur längs der drei Richtungen $B_{ikl} \dot{u}^i \dot{u}^k \dot{u}^l = 0$ fortschreitend drei konsekutive Systemstrahlen antreffen, die 1. ein Paraboloid bestimmen, 2. geodätisch sind. (Man vergleiche hiermit die letzte Konstruktion der Symmetrieflächen im vorigen Paragraphen.)

Es möge noch ein Satz angegeben werden, der sich unmittelbar aus dem Obigen ablesen läßt. Wie im nächsten Paragraphen gezeigt wird, ist in einer linearen Kongruenz:

$$C_{ikl} = \frac{1}{2} B_{ikl} = 0.$$

Daraus folgt nach (2) und (9) der Satz:

Die geodätischen Regelflächen einer linearen Kongruenz sind die parabolischen Flächen zweiten Grades, die der Kongruenz angehören.

Zum Schluß dieses Paragraphen möge noch eine geometrische Interpretation der Form $B_{ikl} d\dot{u}^i d\dot{u}^k d\dot{u}^l = \bar{\Psi}$ angedeutet werden. Dazu berechnen wir die einfachste Invariante der affinen Differentialgeometrie einer Regelfläche unabhängig vom Strahlensystem: Es sei eine Regelfläche gegeben durch den Sechservektor $q(t)$ als Funktionen eines Parameters t , die der Bedingung $q^2 = 0$ genügen und n stetige Ableitungen nach t besitzen, derart, daß nicht alle sechs Komponenten gleichzeitig verschwinden. Man hat die Invarianten einerseits gegen inhaltstreue affine Transformationen des Raumes, andererseits gegen Parametertransformationen und schließlich gegen Umnormierungen der homogenen Plückerschen Linienkoordinaten

zu bestimmen. Setzt man $\frac{dq}{dt} = q'$ und $\frac{d^2q}{dt^2} = q''$, so wird die einfachste Invariante, nach den Regeln des § 2:

$$(10) \quad R^2 = \frac{[q'^2]^{\frac{3}{2}}}{W q q' q'' W}.$$

Man überzeugt sich leicht, daß R einen Rauminhalt darstellt, den man sich durch die zur Fläche $q(t)$ gehörige Schmiege- F_2 konstruieren kann. Gehen wir jetzt wieder in das Strahlensystem und bestimmen die Invariante R^2 für die Regelflächen des Systems, so folgt:

$$(11) \quad R^2 = \frac{(G_{ik} \dot{u}^i \dot{u}^k)^{\frac{3}{2}}}{E_{rs} (\dot{u}^r \ddot{u}^s + \Gamma_{ik}^s \dot{u}^i \dot{u}^k) + B_{ikl} \dot{u}^i \dot{u}^k \dot{u}^l}.$$

Der Nenner ist nichts anderes als die linke Seite der Gleichung (9). Wählt man im besonderen eine geodätische Regelfläche des Strahlensystems, so geht R^2 über in:

$$(12) \quad R^2_{(\text{geod.})} = \frac{(G_{ik} \dot{u}^i \dot{u}^k)^{\frac{3}{2}}}{B_{ikl} \dot{u}^i \dot{u}^k \dot{u}^l} = \frac{\Phi^{\frac{3}{2}}}{\Psi}.$$

Daraus kann man eine Deutung der Form $\bar{\Psi}$ ablesen, da ja die Form Φ durch das Regelement gedeutet ist. Oder: der Quotient $\frac{\Phi^{\frac{3}{2}}}{\bar{\Psi}}$ gibt gerade die niedrigste absolute Invariante der zu der Richtung $\dot{u}^i$ gehörigen geodätischen Regelfläche des Strahlensystems.

§ 9.

Brennflächen.

Im letzten Paragraphen dieses ersten Teiles der affinen Differentialgeometrie der Strahlensysteme soll der Formelapparat auf die *Brennflächen* angewendet werden. Wir begnügen uns damit, die Rechnungen in Torsenparametern auszuführen; also $G_{11} = G_{22} = 0$ vorauszusetzen. Die Brennflächen des Strahlensystems werden gegeben (§ 5) durch die Vektoren.

$$(1) \quad p = y + \frac{1}{4} x \quad \bar{p} = y - \frac{1}{4} x$$

Für die Ableitungsvektoren erhält man nach § 6:

$$(2) \quad \begin{aligned} p_1 &= \frac{1}{4} B'_{i1} x & \bar{p}_1 &= \frac{1}{4} B'_{i1} x - \frac{1}{2} x_1 \\ p_2 &= \frac{1}{4} B'_{i2} x + \frac{1}{2} x_2 & \bar{p}_2 &= \frac{1}{4} B'_{i2} x \end{aligned}$$

Zur Aufstellung der Differentialgleichung der *Asymptotenlinien* der Brennfächen berechnen wir die Determinanten

$$D_{ik} = \begin{vmatrix} p & p_1 & p_2 & p_{ik} \end{vmatrix} \quad \bar{D}_{ik} = \begin{vmatrix} \bar{p} & \bar{p}_1 & \bar{p}_2 & \bar{p}_{ik} \end{vmatrix} \\ = \begin{vmatrix} \left(y + \frac{1}{4}x\right); \frac{1}{2} B_{i1}^l x; \left(\frac{1}{4} B_{i2}^l x + \frac{1}{2} x_2\right); p_{ik} \end{vmatrix} \\ D_{ik} = \frac{1}{8} B_{i1}^l \begin{vmatrix} y & x_1 & x_2 & [p_{ik}]_{x_1} \end{vmatrix} \quad \bar{D}_{ik} = -\frac{1}{8} B_{i2}^l \begin{vmatrix} y & x_1 & x & [\bar{p}_{ik}]_{x_2} \end{vmatrix};$$

dabei sind $[p_{ik}]_{x_1}$ bzw. $[\bar{p}_{ik}]_{x_2}$ die x_1 - bzw. x_2 -Komponente von p_{ik} bzw. $\bar{p}_{ik}$. Dafür erhält man nach § 6:

$$[p_{11}]_{x_1} = \frac{1}{4} B_{i1}^l \quad [\bar{p}_{11}]_{x_2} = -\frac{1}{2} E^{21} B_{111} \\ [p_{12}]_{x_1} = 0 \quad [\bar{p}_{12}]_{x_2} = 0 \\ [p_{22}]_{x_1} = \frac{1}{2} E^{12} B_{222} \quad [\bar{p}_{22}]_{x_2} = \frac{1}{4} B_{i2}^l.$$

also werden:

(3)

$$D_{11} = \frac{1}{8} B_{i1}^l \cdot \begin{vmatrix} y & x & x_2 & x_1 \end{vmatrix} \cdot \frac{1}{4} B_{i1}^l \quad \bar{D}_{11} = -\frac{1}{8} B_{i2}^l \cdot \begin{vmatrix} y & x & x_2 & x_1 \end{vmatrix} \cdot \frac{1}{2} E^{12} B_{111} \\ D_{12} = 0 \quad \bar{D}_{12} = 0 \\ D_{22} = \frac{1}{8} B_{i1}^l \cdot \begin{vmatrix} y & x & x_2 & x_1 \end{vmatrix} \cdot \frac{1}{2} E^{12} B_{222} \quad \bar{D}_{22} = -\frac{1}{8} B_{i2}^l \cdot \begin{vmatrix} y & x & x_2 & x_1 \end{vmatrix} \cdot \frac{1}{4} B_{i2}^l.$$

Und die Differentialgleichung der Asymptotenlinien läßt sich schreiben:

$$(4) \quad \frac{1}{4} B_{i1}^l du^2 + \frac{1}{2} E^{12} B_{222} dv^2 = 0, \quad \frac{1}{2} E^{12} B_{111} du^2 + \frac{1}{4} B_{i2}^l dv^2 = 0.$$

Es ist aber:

$$(5) \quad \frac{1}{4} B_{i1}^l = \frac{1}{4} G^{lr} B_{lr1} = \frac{1}{4} G^{12} B_{121} + \frac{1}{4} G^{21} B_{211} = \frac{1}{2} G^{12} B_{121}$$

und infolge der Normierung $\frac{1}{2} E^{ik} E_{ik} = -1$

$$E^{12} = -G^{12},$$

so daß man für die *Asymptotenlinien der Brennflächen* erhält:

$$(6) \quad B_{121} du^2 - B_{222} dv^2 = 0 \quad \text{und} \quad B_{111}' du^2 - B_{122} dv^2 = 0.$$

Um auch hier auf eine Beziehung zu den Symmetriefflächen hinweisen zu können, führen wir für B_{ikl} den Grundtensor C_{ikl} ein.

Für Torsenparameter ist:

$$(7) \quad \begin{aligned} C_{111} &= -B_{111} & , & \quad C_{122} = -C_{212} = -C_{221} = -B_{122} \\ C_{211} &= -C_{121} = -C_{112} = B_{211}, & \quad C_{222} &= B_{222}. \end{aligned}$$

Dann gehen die Gleichungen (6) über in:

$$(8a, b) \quad C_{112} du^2 + C_{222} dv^2 = 0 \quad \text{und} \quad C_{111} du^2 + C_{221} dv^2 = 0.$$

Daraus erkennt man: wird die erste Gleichung (8) mit dv und die zweite mit du multipliziert und dann die Gleichungen addiert, so erhält man die Differentialgleichung der Symmetriefflächen.

Obwohl wir uns hier auf Torsenparameter beschränkt haben, so mögen dennoch einige wichtige Formeln in allgemeinen Parametern angeführt werden. Bildet man nämlich das Produkt der beiden Gleichungen (8a) und (8b), so erhält man eine biquadratische Gleichung, deren Lösungen die 4 Asymptotenlinien der beiden Brennflächen sind; man erhält:

$$(8c) \quad C_{111} C_{112} du^4 + (C_{111} C_{222} + C_{112} C_{221}) du^2 dv^2 + C_{222} C_{221} dv^4 = 0.$$

Diese Gleichung läßt sich in allgemeinen Parametern in der Form schreiben:

$$(8d) \quad C_{iks} C_{im}^s du^i du^k du^l du^m = 0.$$

Man überzeugt sich leicht, daß (8d) für Torsenparameter bis auf einen unwesentlichen Faktor mit (8c) übereinstimmt.

Die angegebenen Formeln führen sofort zu einigen speziellen Strahlensystemen. Aus den Gleichungen (2) folgt:

1. Ist $B_{121} = 0$ (bzw. $B_{221} = 0$), so wird der Punkt p_1 (bzw. $\bar{p}_2$) unbestimmt, während p_2 und $\bar{p}_1$ bestimmt sind. Das heißt geometrisch: für $B_{121} = 0$ (bzw. $B_{212} = 0$) artet die Brennfläche p (bzw. $\bar{p}$) in eine Kurve aus, deren Tangentenrichtung durch p_2 (bzw. $\bar{p}_1$) gegeben wird. Sind beide Brennflächen des Systems Kurven, so wird $B_{ir}^i = 0$, d. h. die Formen $B_{ikl} du^i du^k du^l$ und $G_{ik} du^i du^k$ sind apolar. Man erkennt aus (7), daß dann auch $C_{ikl} du^i du^k du^l$ zu $G_{ik} du^i du^k$ apolar ist, so daß der Satz gilt: *Sind die quadratische und kubische Grundform des Strahlensystems apolar zueinander,*

dann sind die beiden Brennflächen des Systems in Kurven ausgeartet. Die Bedingungsgleichungen der Apolarität lauten in allgemeinen Parametern:

$$(9) \quad B_{ir}^i = 0 \text{ oder } E^{ik} C_{ikr} = 0.$$

2. Das duale Gegenstück zur Ausartung der Brennflächen des Systems in Kurven ist ihre *Ausartung in Torsen*. Wie aus Gleichung (3) folgt, ist die Brennfläche p (bzw. $\bar{p}$) eine Torse, wenn $B_{222} = 0$ oder $C_{222} = 0$ (bzw. $B_{111} = 0$ oder $C_{111} = 0$). Da die Brennflächen von den Torsen des Strahlensystems längs eines konjugierten Kurvennetzes berührt werden, erkennt man unmittelbar den Satz: Ist eine der beiden Brennflächen eine Torse, dann besteht die eine Schar der Torsen des Systems aus Ebenen, sind beide Brennflächen Torsen, dann sind beide Scharen von Torsen des Systems Ebenen, deren jede eine Brennfläche längs einer Erzeugenden der Brennfläche berührt.

3. Die Ausartungen 1 und 2 treten gleichzeitig auf, wenn $B_{ikl} = 0$ oder $C_{ikl} = 0$; d. h. wenn alle Koeffizienten der kubischen Grundform verschwinden. Dann sind einerseits die beiden Brennflächen des Strahlensystems in Kurven ausgeartet, andererseits sind die Torsen des Systems Ebenen. Das ist aber nur gleichzeitig möglich, wenn jene beiden Kurven gerade Linien sind. *Durch die Gleichungen:*

$$(10) \quad C_{ikl} = 0$$

werden also die linearen Kongruenzen gekennzeichnet. Sie können nicht in ein Geradenbündel bzw. in die Geraden einer Ebene ausarten, da dann unser Formelapparat versagt und wir keinerlei Aussagen machen können (vergl. § 3).

Wesentlich allgemeiner als diese speziellen Fälle sind die

W-Strahlensysteme.

Sie werden dadurch definiert, daß die beiden Brennflächen durch die Strahlen des Systems projektiv aufeinander bezogen werden. Dann müssen die Asymptotenlinien der einen Brennfläche denen der anderen entsprechen. Das besagt nach den Gleichungen (8):

$$(11) \quad C_{111} : C_{221} = C_{112} : C_{222},$$

so daß man die Differentialgleichung der *W-Strahlensysteme* schreiben kann:

$$(11a) \quad \begin{vmatrix} C_{111} & C_{112} \\ C_{221} & C_{222} \end{vmatrix} = 0.$$

Auf den Zusammenhang der kubischen Grundform (Symmetriefflächen) mit den Gleichungen (8) wurde schon hingewiesen. Es möge hier noch eine Eigenschaft im W -Strahlensystem gezeigt werden: die kubische Grundform läßt sich in Torsenparametern in der Form schreiben:

$$(C_{111} du^2 + C_{221} dv^2) du + (C_{112} du^2 + C_{222} dv^2) dv.$$

Im W -Strahlensystem folgt daraus wegen (11) für die Symmetriefflächen:

$$(12) \quad (C_{111} du^2 + C_{221} dv^2) (du + \mu dv) = 0,$$

wo μ einen Proportionalitätsfaktor in (11) andeutet. D. h. geometrisch: Im W -Strahlensystem fallen zwei von den drei Symmetriefflächen durch p mit den gemeinsamen Asymptotenlinien der beiden Brennflächen zusammen, während die dritte aus der besonderen Natur des W -Strahlensystems bestimmt ist. Außerdem sind dann die Asymptotenlinien der Brennflächen zugleich Asymptotenlinien der beiden ausgezeichneten Symmetriefflächen: d. h. zwei der Symmetriefflächen durch p schneiden aus den Brennflächen Asymptotenlinien aus, wie umgekehrt auch die Brennflächen aus den Symmetriefflächen solche ausschneiden. Dies folgt unmittelbar aus den Gleichungen (6), (11), (12) des § 7.

Die Differentialgleichung der W -Strahlensysteme muß sich in invarianter Form ergeben, wenn man den Tensor $C_{iks} C_{e m}^s$ der Gleichung (8d) zweimal mit dem Grundtensor G^{rs} überschiebt. Man erhält dann für die W -Strahlensysteme die außerordentlich einfache Gleichung:

(11b)

$$C^{iks} C_{kis} = 0$$

Der Ausdruck $C^{iks} C_{kis}$ stellt eine simultane Invariante der beiden Grundformen des Strahlensystems dar. Es bleibt einer späteren Mitteilung vorbehalten, die simultanen Invarianten systematisch zu untersuchen.

(Eingegangen: 23. VII. 1928.)

Théorème sur trois continus.

Par Casimir Kuratowski à Lwów.

Je vais prouver dans cette note le théorème suivant:

Théorème. A , B et C étant trois continus situés sur le plan et tels que $A + B + C$ est une coupure¹⁾ entre deux points p et q , tandis que ni $A + B$, ni $A + C$, ni $B + C$ n'est une coupure entre ces points, on a $ABC = 0$ ²⁾.

1. J'aurai recours dans la démonstration à un théorème de Janiszewski et à deux lemmes que j'établirai tout-à-l'heure.

D'après le théorème de Janiszewski [généralisé³⁾], si X et Y sont deux ensembles fermés dont un au moins est borné et dont le produit est un continu, si ni X ni Y n'est un $\mathfrak{S}_{(p, q)}$, $X + Y$ n'en est un nonplus.

Lemme 1. F et D étant deux ensembles fermés et bornés, dont D est un continu, tels que $F + D$ est un $\mathfrak{S}_{(p, q)}$, l'ensemble $F - D$ contient un ensemble connexe S tel que $D + S$ est un $\mathfrak{S}_{(p, q)}$.

Démonstration. La propriété „d'être un ensemble fermé qui est un $\mathfrak{S}_{(p, q)}$ contenant D étant inductive⁴⁾, $F + D$ contient un ensemble J irréductible par rapport à cette propriété. L'ensemble $S = J - D$ est l'ensemble cherché.

En effet, dans le cas contraire, on aurait $S = M + N$, M et N étant deux ensembles séparés non-vides. Par conséquent, les ensembles $D + M$ et $D + N$ sont deux ensembles fermés satisfaisant aux formules:

¹⁾ Un ensemble X est dit une „coupure entre p et q “, ou en symboles: un $\mathfrak{S}_{(p, q)}$, lorsque tout continu qui unit p et q passe par X et les points p et q sont situés en dehors de X .

²⁾ Pour le cas, où $A + B + C$ est une frontière commune à deux régions, ce théorème fut prouvé par M. Knaster; je m'en suis servi (de ce cas particulier) dans la démonstration du théorème „fondamental“ de ma note „Sur la structure des frontières communes à deux régions“, Fund. Math., XII, p. 29.

³⁾ Prace mat.-fiz., 26 (1913).

⁴⁾ C'est-à-dire qu'elle appartient au produit de toute suite d'ensembles décroissants jouissants de cette propriété. D'après un théorème de M. Brouwer, tout ensemble jouissant d'une propriété inductive qui implique que l'ensemble est fermé contient un ensemble irréductible par rapport à cette propriété (c'est-à-dire, que ce dernier ensemble jouit de cette propriété tandis qu'aucun de ses vrais sous-ensembles ne la possède).

$$(1) \quad D + M \neq J \neq D + N$$

$$(2) \quad (D + M)(D + N) = D.$$

Selon (1), ni $D + M$, ni $D + N$ n'est un $\mathfrak{S}_{(p, q)}$, donc, selon (2) et le théorème de Janiszewski, $(D + M) + (D + N)$ n'en est un nonplus. Mais ceci est impossible, car $(D + M) + (D + N) = J$ et J est un $\mathfrak{S}_{(p, q)}$.

Lemme 2. A , B et C étant trois continus bornés tels que 1°: $A + B + C$ est un $\mathfrak{S}_{(p, q)}$, 2°: ni $A + C$, ni $B + C$ n'est un $\mathfrak{S}_{(p, q)}$, — il existe, pour chaque point a de C , un arc simple ax tel qu'on peut remplacer C par ax dans 1° et 2°.

Démonstration. On peut évidemment supposer que

$$(3) \quad A + B \text{ n'est pas un } \mathfrak{S}_{(p, q)},$$

car, en cas contraire, ax est un arc arbitraire suffisamment petit.

En vertu de 2°, il existe deux continus L et M qui unissent p et q et que

$$(4) \quad L(A + C) = 0 = M(B + C).$$

Il résulte de là que $(L + M)C = 0$; il existe, par conséquent, une courbe simple fermée (même une ligne polygonale) F qui coupe le plan entre les continus $L + M$ et C . Soit R celle des deux régions en lesquelles F coupe le plan qui contient C . On a donc:

$$(5) \quad C \subset R$$

$$(6) \quad (L + M)\bar{R} = 0^5),$$

d'où, selon (4): $L(A + \bar{R}) = 0 = M(B + \bar{R})$, ce qui prouve que

$$(7) \quad \text{ni } A + \bar{R} \text{ ni } B + \bar{R} \text{ n'est un } \mathfrak{S}_{(p, q)}.$$

Je dis que

$$(8) \quad A + B + F \text{ est un } \mathfrak{S}_{(p, q)}.$$

Soit, en effet, X un continu qui unit p à q . Si $X(A + B) = 0$, on conclut de 1° que $XC \neq 0$, d'où en raison de (5): $XR \neq 0$. Or, X , comme continu ayant des points communs avec R ainsi qu'avec le complémentaire de R (puisque R ne contient pas p), X a des points communs avec la frontière de R , c'est-à-dire avec F .

La proposition (8) est ainsi établie.

En vertu du théorème de Janiszewski, on conclut de 1° et 2° que le produit $(A + C)(B + C)$ n'est pas un continu, donc que $AB \neq 0$. Il résulte de là que $A + B$ est un continu. On peut donc,

⁵⁾ $\bar{R}$ désigne l'ensemble composé de R et de ses points limites.

en tenant compte de (8), substituer, dans le lemme 1, $A + B$ à la place de D .

On en conclut l'existence d'un ensemble connexe S tel que $S \subset F - (A + B)$ et que

$$(9) \quad A + B + S \text{ est un } \mathfrak{S}_{(p, q)}.$$

L'ensemble F , comme frontière de la région R , est contenu dans $\bar{R}$, donc, selon (7), F n'est pas un $\mathfrak{S}_{(p, q)}$. Il en résulte, en vertu de (3) et (8), que le produit $(A + B)F$ n'est pas un continu; il ne peut donc être ni vide ni composé d'un seul point. Par conséquent, S diffère de la courbe simple fermée F de deux points au moins. Il existe donc un arc xy de cette courbe tel que

$$(10) \quad S \subset xy.$$

Unissons a à y par un arc ay situé dans $R + y$. Soit ax l'arc composé de ay et yx . On a donc $ax < \bar{R}$, d'où selon (7), ni $A + ax$, ni $B + ax$ n'est un $\mathfrak{S}_{(p, q)}$; d'autre part, $A + B + ax$ est un $\mathfrak{S}_{(p, q)}$, en raison de (9) et (10).

2. Démonstration du théorème.

Nous allons démontrer d'abord le théorème pour le cas où A , B et C sont bornés.

Supposons, contrairement à la thèse du théorème, qu'il existe un point a tel que

$$(11) \quad a \subset ABC.$$

Selon le lemme 2, il existe un arc ax tel que:

$$(12) \quad A + B + ax \text{ est un } \mathfrak{S}_{(p, q)}$$

$$(13) \quad \text{ni } A + ax, \text{ ni } B + ax \text{ n'est un } \mathfrak{S}_{(p, q)}.$$

Soit b le premier point de l'arc ax (rangé de a à x) tel que

$$(14) \quad A + B + ab \text{ est un } \mathfrak{S}_{(p, q)}.$$

Un tel point existe, car la propriété d'être un arc ax assujetti à la condition (12) est inductive.

En posant dans le lemme 1: $F = ab$ et $D = A + B$, on en conclut que l'ensemble $ab - (A + B)$ contient un ensemble connexe S tel que $A + B + S$ est un $\mathfrak{S}_{(p, q)}$. Soit $a_1 b_1$ le plus petit arc extrait de ab qui contient S ; on a donc

$$a \leq a_1 < b_1 \leq b.$$

De plus: $b_1 = b$, car, autrement, l'ensemble $A + B + ab_1$ ne serait pas (conformément à la définition de b) un $\mathfrak{S}_{(p, q)}$, tandis que $A + B + S \subset A + B + ab_1$, et $A + B + S$ est un $\mathfrak{S}_{(p, q)}$.

Or, soit d un point intérieur de l'arc a_1b . Donc

$$(15) \quad A + B + ad \text{ n'est pas un } \mathfrak{S}_{(p, q)}.$$

D'autre part, selon la définition de l'arc a_1b_1 , on a $db \subset S + b$ et, comme $S(A + B) = 0$, il vient:

$$(16) \quad db \cdot (A + B) \subset b,$$

ce qui prouve que l'ensemble $db \cdot (A + B)$ se réduit à un seul point (à moins qu'il ne soit vide), il est donc contenu soit dans A , soit dans B . Par raison de symétrie, on peut supposer que

$$(17) \quad db \cdot (A + B) \subset A.$$

Or, considérons les deux ensembles: $A + B + ad$ et $A + ab$. On a: $(A + B + ad)(A + ab) = (A + ad + B)(A + ad + db) = A + ad + B \cdot db$,

ce qui, selon (17), est égal à $A + ad$ et, comme d'après (11), a appartient à A , $A + ad$ est un continu.

En appliquant aux ensembles $A + B + ad$ et $A + ab$ le théorème de Janiszewski, en tenant compte de (13) et (15), on en conclut que leur somme, c'est-à-dire, l'ensemble $A + B + ab$ n'est pas un $\mathfrak{S}_{(p, q)}$. Mais ceci contredit la proposition (14).

Notre théorème est ainsi établi dans l'hypothèse que A , B et C sont bornés.

Dans le cas où cette hypothèse n'est pas réalisée, on transforme le plan par une inversion ayant pour pôle un point situé en dehors de $A + B + C + p + q$. En tenant compte du fait que la propriété d'être un $\mathfrak{S}_{(p, q)}$ fermé reste invariante lorsque on ajoute le pôle à l'ensemble transformé⁶⁾, on est ramené au cas où A , B et C sont bornés.

On voit, en même temps, que lorsque tous les trois ensembles A , B et C sont non-bornés, leurs images ont toujours le pôle pour point d'accumulation, même dans le cas où $ABC = 0$. On arrive ainsi au corollaire suivant, dû à M. Knaster:

Corollaire. A , B et C étant trois continus non-bornés situés sur le plan, si $A + B + C$ est un $\mathfrak{S}_{(p, q)}$, l'un des trois ensembles $A + B$, $A + C$ ou $B + C$ en est un aussi.

⁶⁾ Cf. Fund. Math., V, p. 32 (note de M. Knaster et moi).

Concerning upper semi-continuous collections.

By R. L. Moore (Austin, Texas, U. S. A.).

I have shown¹⁾ that if, in a plane S , G is an upper semi-continuous collection of mutually exclusive compact continua and no continuum of the collection G separates S and every point of S belongs to some continuum of G then the collection G is, with reference to a suitably defined notion of limit element, topologically equivalent to a plane. It follows, in an obvious manner, that the same proposition holds true if the word plane is replaced throughout by sphere. In the present paper I will consider the situation obtained on the removal of the restriction that no continuum of the set G shall separate S . For the sake of simplicity certain of the results will be established for the case where S is a sphere instead of a plane. From these results corresponding results for the plane may be readily obtained.

Definition. A cactoid (opuntioid) is a bounded continuous curve M lying in space of three dimensions and such that (a) every non-degenerate maximal cyclic subset²⁾ of M is a simple closed surface and (b) no point of M lies in a bounded complementary domain of any subcontinuum of M . A cactoid M is said to be aspiculate if it does not contain any arc t such that no point

¹⁾ Concerning upper semi-continuous collections of continua, Transactions of the American Mathematical Society, vol. 27 (1925), pp. 416—428.

²⁾ A cyclic subset of a continuous curve M is a subset K of M such that every two points of K belong to some simple closed curve which is a subset of K and a cyclic subset of M is said to be maximal if it is not a proper subset of any other cyclic subset of M . See the following papers by G. T. Whyburn: Cyclically connected continuous curves, Proceedings of the National Academy of Sciences, vol. 13 (1927), pp. 31—38; Some properties of continuous curves, Bulletin of the American Mathematical Society, vol. 33 (1927), pp. 305—308; Concerning the structure of a continuous curve, American Journal of Mathematics, vol. L (1928), pp. 167—194. The term maximal cyclic subset as used here is more inclusive than Whyburn's maximal cyclic continuous curve only in that a point may be a maximal cyclic subset of M but no point is a cyclic continuous curve in the sense of Whyburn. On the other hand Whyburn's cyclic element is more inclusive than maximal cyclic subset in that a cut point of M that belongs to a maximal cyclic continuous curve of M is a cyclic element of M but not a maximal cyclic subset of M . This slight departure from Whyburn's terminology is made, in the present paper, largely for contextual reasons. A maximal cyclic subset of M will be said to be degenerate if it consists of a single point.

of t except its endpoints is a limit point of $M-t$. An aspiculate cactoid M is simple if no two cut points of M are separated in M by infinitely many different points.

Theorem 1. If G is an upper semi-continuous collection of mutually exclusive continua filling up a spherical surface S and A , B and O are three different continua of the set G and there exists a continuum of G that separates A from O , but no continuum of G separates A from both O and B , then there exists a continuum g of G such that (a) g separates A from O but (b) no continuum of G separates A from g .

Theorem 2³). If G is an upper semi-continuous collection of mutually exclusive continua filling up a spherical surface S then the space S_G whose elements are the continua of G is topologically equivalent to a cactoid.

Proof. If the elements of S_G are called points and the term region is defined as on page 428 of my paper Concerning upper semi-continuous collections of continua then⁴) Axioms 1, 2, 4 and 5₁ and Theorem 4 of my paper On the foundations of plane analysis situs⁵) all hold true and an element p of S_G is a limit element of a set H of elements of S_G if and only if every region that contains p contains at least one element of H distinct from p . It follows that S_G is a separable metric space. Clearly S_G is compact. Suppose now that Q is a non-degenerate maximal cyclic subset of the space S_G . The set of elements Q contains a simple closed curve of elements of S_G and no such curve contains⁶) more than a

³) Vietoris has shown that if the space S_G is one dimensional then it is an acyclic continuous curve. Every such curve is a cactoid whose maximal cyclic subsets are all points. Cf. L. Vietoris, Über stetige Abbildungen einer Kugelfläche, Proceedings of the Royal Academy of Sciences of Amsterdam, vol. 29 (1926), pp. 443—453.

⁴) R. L. Moore, loc. cit., Theorem 26. While this theorem was established for a Euclidean space S of n -dimensions, it is clear that (if the last definition of Page 417 is omitted) the same argument suffices to establish it for the case where the space S is a spherical surface or, indeed, any metric space which itself satisfies Axioms 1, 2, 4 and 5, and Theorem 4 of my Foundations article. That every space satisfying Axiom 1 and Theorem 4 is necessarily metric may be easily seen with the help of the Urysohn-Tychonoff theorem to the effect that every regular and perfectly separable Hausdorff space is metric. See P. Urysohn, Mathematische Annalen, vol. 94 (1925), pp. 309—315 and A. Tychonoff, Ibid. vol. 95 (1926), pp. 139—142. It has been shown by Chittenden that in order that a topological space in which there are no isolated points should be metric and separable it is necessary and sufficient that it should satisfy my Axiom 1. Cf. E. W. Chittenden, On the metrization problem and related problems in the theory of abstract sets, Bulletin of the American Mathematical Society, vol. 33 (1927), pp. 13—34.

⁵) Transactions of the American Mathematical Society, vol. 17 (1916), pp. 131—164.

⁶) See S. Mazurkiewicz, Fundamenta Mathematicae, vol. 2 (1921), pp. 119—130 and R. L. Moore, Proceedings of the National Academy of Sciences, vol. 9 (1923), pp. 101—106.

countable number of cut elements of S_G . Thus there exists a continuum A of G belonging to the set Q and not separating S and such that not every continuum of G distinct from A is separated from A by some continuum of G . If no continuum of G separates S then S_G is equivalent to a sphere. Suppose there exists at least one continuum of G that separates S . Let H denote the set of continua consisting of A together with all continua y of G such that there exists no continuum of G separating y from A . By Theorem 1, if g is a continuum of G distinct from A and there exists a continuum of G that separates g from A then there exists a continuum X_g belonging to H and such that (a) X_g separates g from A , (b) no continuum of G separates A from X_g . The collection H is an upper semi-continuous collection of continua. The elements of H form a simple closed surface S_H of elements⁷⁾. This set of elements is identical with Q . For suppose, on the contrary, that there exists a simple closed curve J of elements of S_G that contains two elements h_1 and h_2 of H and an element g that does not belong to H . The continuum X_g separates g from one of the continua h_1 and h_2 and therefore there exists no simple closed curve of elements of G containing g , h_1 and h_2 . Thus the supposition that H is not identical with Q leads to a contradiction. Hence every non-degenerate maximal cyclic subset Q of S_G is a simple closed surface of elements. That S_G is homeomorphic with a cactoid may now be established by a modification of arguments that have been given⁸⁾ to show that every acyclic continuous curve lying in n -dimensional space is homeomorphic with some plane acyclic continuous curve.

Theorem 3. If, on a sphere S , $M_1, M_2, M_3 \dots$ is a sequence of continua, all proper subsets of S , and $P_1, P_2, P_3 \dots$ is a sequence of distinct points then there exists on S a contracting⁹⁾ sequence of mutually exclusive continua $N_1, N_2, N_3 \dots$ such that (a) for each n , there exists a continuous transformation of S into itself that throws M_n into N_n , (b) if the continua of the sequence $N_1, N_2, N_3 \dots$ are regarded as elements and each point that neither belongs to a continuum of this sequence nor is separated from one continuum of this sequence by another one is regarded as an element then between the set of elements so obtained and the set of all points of S there exists a continuous one to one correspondence in which, for each n , N_n corresponds to P_n .

⁷⁾ Cf. L. Vietoris, loc. cit., page 7, fourth paragraph.

⁸⁾ See T. Wazewski, Sur les courbes de Jordan ne renfermant aucune courbe simple fermée de Jordan, Annales de la Société Polonaise de Mathématique, vol. 2 (1923), pp. 49—170, and a later article, by H. M. Gehman, Transactions of the American Mathematical Society, vol. 29 (1927), pp. 553 bis 568.

⁹⁾ A sequence of point sets is said to be contracting if, for each positive number ϵ , all but a finite number of the point sets of that sequence are of diameter less than ϵ .

Theorem 4. If t is a simple chain¹⁰⁾ of maximal cyclic subsets of a cactoid K , M is the sum of all the elements of t and S is a sphere, then there exists, on S , an upper semi-continuous collection G of mutually exclusive continua and a one to one reciprocal continuous correspondence between the points of M and the continua of G such that if, for each point x of M , $T(x)$ denotes the corresponding continuum of G then, for each x , the number of complementary domains of $T(x)$ on the sphere S equals the number of components of $K-x$ and this set of domains, if infinite, forms a contracting sequence.

Proof. Let r denote some definite simple continuous arc lying on S . Let H denote the point set consisting of all the cut points of M together with each point (if there is any) which is an end-element of t . Let $P_1, P_2, P_3 \dots$ denote the set¹¹⁾ of all points X of H such that either (1) $K-X$ has more than two components or (2) X is an end-element of t and $K-X$ has more than one component. There exists, on r , a closed point set V which is in such a one to one reciprocal correspondence with H that if, for each point X of H , $W(X)$ denotes the point of V associated with X in this correspondence then $W(X)$ separates $W(Y)$ from $W(Z)$ on r if, and only if, X separates Y from Z in M . There exists, on S , a sequence of mutually exclusive continua $L_1, L_2, L_3 \dots$ with the same number of terms (finite or countably infinite) as the sequence $P_1, P_2, P_3 \dots$ and such that (1) the diameter of L_i is equal to or less than $1/i$, (2) L_i has, in common with r , only the point $W(P_i)$, (3) L_i is the sum of a finite or countably infinite number of simple closed curves forming a contracting sequence Q_i , (4) Q_i is an infinite sequence if $K-P_i$ has infinitely many components and if the number of components of $K-P_i$ is finite and equal to n_i then the number of simple closed curves in Q_i is n_i-1 or n_i-2 according as P_i is or is not an end-element of t , (5) no point of any curve of the sequence Q_i is separated, on S , from any point of r by any other curve of this sequence and no two points of r are separated from each other, on S , by any curve of Q_i , (6) every two curves of Q_i have in common only the point $W(P_i)$.

If every continuum of the sequence $L_1, L_2, L_3 \dots$ is regarded as an element and every point of S which neither belongs to any

¹⁰⁾ A simple chain of maximal cyclic subsets of a cactoid K is a collection t of maximal cyclic subsets of K such that if M denotes the sum of all the elements of t then M is a continuum containing two elements A and B of t such that A and B do not both belong to any proper connected subset of M which is the sum of the elements of some collection of maximal cyclic subsets of K . Cf. G. T. Whyburn, Concerning the structure of a continuous curve, loc. cit., page 169.

¹¹⁾ That the set of all such points is countable may be established by an argument similar to arguments employed to show that no acyclic continuous curve has more than a countable number of branch points. Cf. T. Wazewski, loc. cit. and K. Menger, Über reguläre Baumkurven, Mathematische Annalen, vol. 96 (1926), pp. 572-582.

continuum of this sequence nor is separated from any point of r by any continuum of this sequence is regarded as an element, the set U of elements so obtained is upper semi-continuous. It follows that between this set of elements and the set of all points of a sphere I there exists a one to one continuous correspondence in which the set of elements of U consisting of $L_1, L_2, L_3 \dots$ and all points of r that belong to no continuum of this sequence corresponds to a simple continuous arc z of points on I , z being a great semicircle or a proper subinterval of a great semicircle according as it is or is not true that both endelements of t are points. For each element x of U let $W_1(x)$ denote the point of I that is associated with x in this correspondence. For each point X of the set $H - (P_1 + P_2 + P_3 + \dots)$ let $V(X)$ denote $W_1 W(X)$. Let $V(P_i)$ denote $W_1(L_i)$. If at least one of the endelements of t is a point let F denote one such point and let O denote $V(F)$. If neither endelement of t is a point let O denote a point not lying on the arc z and not opposite to either endpoint of z but lying on a great semicircle of which z is a part. For each point X of the set H let $C(X)$ denote the circle (degenerate or not), with center at O , that lies on I and passes through $V(X)$, let $U(X)$ denote the collection of all continua Y of U such that $W_1(Y)$ is a point of $C(X)$ and let $Q(X)$ denote the point set obtained by adding together all the continua¹²⁾ of $U(X)$. Let U_t denote the collection of all such continua $Q(X)$ for all points X of H . Let N denote the closed point set obtained by adding together all the continua of the collection U_t . Let $S_1, S_2, S_3 \dots$ denote the simple closed surfaces, if there are any, which are elements of t . If S_n is not an endelement of t it contains two points X_n and Y_n which are cut points of M . If it is an endelement of t then S_n contains one, and only one, point X_n which is a cut point of M . On the sphere S let D_n denote a complementary domain of N whose boundary contains every element of $U(X_n)$ which is a point or every element of $U(X_n) + U(Y_n)$ which is a point, according as S_n is or is not an endelement of t . Suppose S_n contains one or more cut points of K that are not cut points of M . There are¹³⁾ only a countable number of such points. Call them $P_{n1}, P_{n2}, P_{n3} \dots$. It may be shown that there exists a contracting sequence of mutually exclusive continua $E_{n1}, E_{n2}, E_{n3} \dots$, all lying in D_n , such that (a) the sequences $E_{n1}, E_{n2}, E_{n3} \dots$ and $P_{n1}, P_{n2}, P_{n3} \dots$ have the same number of terms, finite or countably infinite, (b) the number of complementary domains of E_{nj} on S equals the number of components of $K - P_{nj}$, (c) the point set $(N + E_{n1} + E_{n2} + E_{n3} \dots) - E_{nj}$ lies in one complementary domain of E_{nj} on S (that is to say in one component of $S - E_{nj}$), (d) if R denotes the collection whose elements are $Q(X_n), Q(Y_n)$ in case Y_n exists, $E_{n1}, E_{n2}, E_{n3} \dots$ and the points of D_n that are not separated from N on S by any

¹²⁾ For each X , all but possibly one of the continua of $U(X)$ are degenerate.

¹³⁾ See S. Mazurkiewicz, loc. cit. and R. L. Moore, loc. cit.

E_{nj} then between the elements of the upper semi-continuous collection R and the points of S_n there exists a one to one continuous correspondence Z_n such that if, for each point X of S_n , $Z_n(X)$ denotes the corresponding element of R then $Q(X_n) = Z_n(X_n)$, for each P_{nj} , $E_{nj} = Z_n(P_{nj})$ and, if Y_n exists, $Q(Y_n) = Z_n(Y_n)$. Let T denote a transformation of the points of M into continua on S such that (1) if the point X belongs to H then $T(X) = Q(X)$, (2) if X is a point which does not belong to H but which lies on the simple closed surface S_n then $T(X) = Z_n(X)$. The transformation T satisfies all the conditions required of T in Theorem 4.

Theorem 5.¹⁴) If K is a cactoid and S is a sphere there exists on S an upper semi-continuous collection G of continua such that the space whose elements are the continua of G is topologically equivalent to K and such that every nondegenerate continuum of G separates S .

Proof. There exists¹⁵) a finite or countably infinite contracting sequence $t_1, t_2, t_3 \dots$ of simple chains of maximal cyclic subsets of K such that (a) neither endelement of t_1 separates K , (b) for each n greater than 1, one endelement of t_n is an element of t_m for some m less than n , (c) if $i \neq j$ and t_i and t_j have an element in common that element is an endelement either of t_i or of t_j , (d) if there exists a maximal cyclic subset of K that is not an element of any t_n , every such subset of K is a point and the set of all such points is totally disconnected. For each n let M_n denote the sum of all the elements of t_n . Let G_1 and T_1 respectively denote an upper semi-continuous collection of continua and a continuous correspondence satisfying with reference to t_1 the conditions which, according to the statement of Theorem 4, are satisfied by G and T with reference to t , and satisfying the additional condition that if X is a cut point of K belonging to M_1 then every component of $S - T_1(X)$ that contains no point of any element of G_1 is of diameter less than 1. Let E_1 denote the sum of all the continua of G_1 and if X is a cut point of K belonging to M_1 let $D_{x_1}, D_{x_2}, D_{x_3} \dots$ denote the components of $S - T_1(X)$, if there are any, that contain no points of E_1 . Let A_2 denote that endelement of t_2 which belongs to t_1 . If A_2 is a point let X_2 denote A_2 . If A_2 is a simple closed surface let X_2 denote the point of this surface which is a cut point of t_2 . In either case there exists a collection G_2 of mutually exclusive continua and a correspondence T_2 satisfying with reference to t_2 the conditions which, according to the statement of Theorem 4, are satisfied by G and T with reference to t and such that if E_2 denotes the sum of all the continua of G_2 then (a) if X is a point of A_2 , $T_2(X) = T_1(X)$, (b) if X is a point not belonging to A_2 but belonging to some element

¹⁴) For the case where K is an acyclic continuous curve see L. Vietoris, loc. cit.

¹⁵) The truth of this statement may be established by an argument having much in common with the argument given by R. L. Wilder to prove Theorem 15 of his paper Concerning continuous curves, *Fundamenta Mathematica*, vol. 7 (1925), pp. 340—377.

of t_2 then $T_2(X)$ lies in the domain D_{X_2} , and (c) if X is a cut point of M_2 , not belonging to A_2 , then every component of $S - T_2(X)$ that contains no point of E_2 is of diameter less than $1/2$. If X is a cut point of K belonging to M_2 and not to M_1 let D_{X_1} , D_{X_2} , D_{X_3} . . . denote the components of $S - T_2(X)$, if there are any, that contain no points of E_2 . Let A_3 denote the endelement of t_3 which belongs to t_1 or to t_2 . Let X_3 denote A_3 or the cut point of M_3 that lies on A_3 according as A_3 is a point or a surface and, in either case, let G_3 and T_3 denote a collection of continua and a correspondence satisfying with reference to t_3 the conditions which, according to Theorem 4, are fulfilled by G and T with reference to t and such that if E_3 denotes the sum of all the continua of G_3 then (a) if X is a point of A_3 , $T_3(X) = T_1(X)$ or $T_2(X)$ according as X_3 is or is not in M_1 , (b) if X is a point not belonging to A_3 but belonging to some element of t_3 then $T_3(X)$ lies in the domain D_{X_2} , or the domain D_{X_3} , according as X_3 is or is not identical with X_2 and (c) if X is a cut point of M_3 , not belonging to A_3 , then every component of $S - T_3(X)$ that contains no point of E_3 is of diameter less than $1/3$. This process may be continued. Thus, in case t_1, t_2, t_3 . . . is an infinite sequence, we have an infinite sequence G_1, G_2, G_3 . . . of collections of continua and a sequence T_1, T_2, T_3 . . . of correspondences such that, for each n greater than 1, G_n and T_n satisfy with respect to t_n the conditions which, according to the statement of Theorem 4, are satisfied by G and T with reference to t and such that if E_n denotes the sum of all the continua of G_n and for each cut point X of K belonging to M_n but not to M_j for any j less than n , $D_{X_1}, D_{X_2}, D_{X_3}$. . . denote the components of $S - T_n(X)$ that contain no points of E_n and m denotes the smallest integer such that t_m contains an endelement of t_n and A_n denotes the endelement of t_n that belongs to t_m and X_n denotes A_n or the cut point of t_n that lies on A_n according as A_n is a point or a surface, then (a) if X is a point of A_n , $T_n(X) = T_m(X)$ and (b) if X is a point of $M_n - A_n$ then $T_n(X)$ lies in the first domain of the sequence $D_{X_{n1}}, D_{X_{n2}}, D_{X_{n3}}$. . . that does not contain $T_j(X)$ for any point X that belongs to M_j for some j less than n and (c) if X is a cut point of $M_n - A_n$ then every component of $S - T_n(X)$ that contains no point of E_n is of diameter less than $1/n$.

If X is a point of M_n let $T(X)$ denote the continuum $T_n(X)$. If X is a point of K belonging to no M_n there exists a sequence of integers n_1, n_2, n_3 . . . such that if, for each j , P_{n_j} is a point belonging to M_{n_j} then X is the sequential limit point of the sequence $P_{n_1}, P_{n_2}, P_{n_3}$. . . It may be seen that the sequence of continua $T(P_{n_1}), T(P_{n_2}), T(P_{n_3})$. . . has a sequential limit point and that if $P_{m_1}, P_{m_2}, P_{m_3}$. . . is another sequence of points such that, for each j , P_{m_j} belongs to M_{m_j} and such that X is the sequential limit point of $P_{m_1}, P_{m_2}, P_{m_3}$. . . then $T(P_{m_1}), T(P_{m_2}),$

$T(P_{m_3}) \dots$ and $T(P_{n_1}), T(P_{n_2}), T(P_{n_3}) \dots$ have the same sequential limit point. Call this limit point $T(X)$. For every point of K there has now been defined a corresponding element $T(X)$ of S . The set of all such elements $T(X)$ is upper semi-continuous and the transformation T is continuous.

Theorem 6. If, on a sphere S , H is a contracting sequence of mutually exclusive continua and G is the collection whose elements are the continua of the sequence H and the points of S that belong to no continuum of H (in other words if G is a collection of continua which fills up S and the nondegenerate elements of G form a contracting sequence) then the space S_G whose elements are the continua of G is topologically equivalent to a simple aspiculate cactoid.

That the collection G is upper semi-continuous is clear. It follows, by Theorem 2, that S_G is topologically equivalent to a cactoid. It is easy to see that every such cactoid is a simple aspiculate one.

Theorem 7. If K is a simple aspiculate cactoid and S is a sphere there exists on S an upper semi-continuous collection G of mutually exclusive continua such that the space whose elements are the continua of G is topologically equivalent to K and such that the non-degenerate continua of G form a contracting sequence.

The truth of Theorem 7 may be seen with the aid of results established in the above proof of Theorem 5.

It has been shown by J. R. Kline¹⁶⁾ that if, in a plane, H is a contracting sequence of mutually exclusive compact continua, the sum of all the continua of H is not connected. It may be seen that the argument given to prove Theorem 26 of my paper Concerning upper semi-continuous collections of continua¹⁷⁾ holds true throughout if the word „continuum“ is replaced by „closed point set“. It follows that if H is a contracting sequence of mutually exclusive closed and compact point sets in a space S of n -dimensions and G is the collection whose elements are the point sets of H and the points of S that belong to no point set of H , then the term region may be so defined that (1) if the elements of G are called points, Axioms 1, 2, 4 and 5₁ and Theorem 4 of my Foundations paper all hold true and, furthermore, an element p of G will be a limit element of a set W of elements of G if and only if every region that contains p contains at least one element of W distinct from p . It follows that the space whose points are the elements of G is metric. But, in a metric space, no countable point set is connected¹⁸⁾. Therefore Kline's theorem holds true for closed point sets as well as for continua.

¹⁶⁾ J. R. Kline, Concerning the sum of a countable infinity of mutually exclusive continua, *Mathematische Zeitschrift*, vol. 26 (1927), pp. 687—690.

¹⁷⁾ Loc. cit., Section IV, pages 427 and 428.

¹⁸⁾ Cf. Hausdorff, *Grundzüge der Mengenlehre*, page 248.

Über affektfreie Gleichungen.

Von Ph. Furtwängler in Wien.

Im folgenden soll eine ziemlich allgemeine Gestalt affektfreier Gleichungen angegeben werden¹⁾, die so übersichtlich ist, daß sie gestattet, die Richtigkeit des folgenden Satzes, dessen Beweis das Hauptziel dieser Entwicklungen bildet, festzustellen. Ist eine Gleichung n^{ten} Grades mit reellen Koeffizienten und dem ersten Koeffizienten 1

$$x^n + a_1 x^{n-1} + \dots + a_{n-1} x + a_n = 0 \quad (1)$$

vorgelegt und interpretiert man die Koeffizienten a_i als rechtwinklige Koordinaten eines R_n , so entspricht jeder Gleichung eindeutig ein Punkt des Koeffizientenraumes R_n und umgekehrt. Mit dieser Bezeichnungsweise lautet der zu beweisende Satz:

Die Gleichungen n^{ten} Grades mit Koeffizienten aus einem beliebig vorgegebenen algebraischen reellen Rationalitätsbereich k und dem ersten Koeffizienten 1, die in k keinen Affekt haben, liegen im Koeffizientenraum R_n überall dicht.

Aus dem Satz folgt z. B., daß man die Anzahl der reellen Wurzeln $r \equiv n(2)$ für die affektfreien Gleichungen beliebig vorschreiben kann. Ein analoger Satz läßt sich aufstellen, wenn der zu Grunde gelegte algebraische Rationalitätsbereich komplex ist.

Ich bediene mich bei den folgenden Entwicklungen der Newtonschen Polygone (*N. P.*), die für Fragen der Irreduzibilität zuerst von G. Dumas²⁾ herangezogen sind; später ist ihre Verwendung von Ö. Ore³⁾ ausführlich und systematisch in zahlreichen Arbeiten studiert. Wir brauchen hier nur einen kleinen Ausschnitt dieser umfassenden Theorie, den ich in § 1 zusammenstelle und auf möglichst kurzem Wege begründe. Ich setze nur die Grundbegriffe der Idealtheorie in algebraischen Zahlkörpern als bekannt voraus.

¹⁾ Mit dem Problem der Aufstellung affektfreier Gleichungen auf algebraischem Wege beschäftigt sich eine ganze Reihe von Abhandlungen. Man findet ausführlichere Literaturangaben bei O. Perron, Über Gleichungen ohne Affekt, Heidelberg. Akad. Sitz. Ber. 1923.

²⁾ J. de math. 6. sér. t. 2 (1906), p. 191.

³⁾ Man vergl. besonders Acta math. 44 (1923), p. 219 und Math. Zeitschr. 18 (1923), p. 278.

§ 1.

Ich benutze folgende bekannte Ausdrucksweise: Eine algebraische Zahl α hat bezüglich eines Primideals $\mathfrak{P}$, das in $k(\alpha)$ oder einem Oberkörper von $k(\alpha)$ liegt, die (ganzzahlige) Ordnung e , wenn $\frac{\alpha}{\mathfrak{P}^e} \bmod. \mathfrak{P}$ ganz, dagegen $\frac{\alpha}{\mathfrak{P}^{e+1}} \bmod. \mathfrak{P}$ nicht ganz ist. Ich werde in diesem Falle auch sagen, daß α genau durch $\mathfrak{P}^e$ teilbar ist.

Es sei nun

$$f(x) = x^n + \alpha_1 x^{n-1} + \dots + \alpha_{n-1} x + \alpha_n, \quad \alpha_n \neq 0 \quad (2)$$

ein Polynom mit Koeffizienten aus einem Zahlkörper k und $\mathfrak{P}$ ein Primideal aus k oder einem Oberkörper von k . Besitzt dann der Koeffizient α_i bez. $\mathfrak{P}$ die Ordnung e_i , so markiere man in einem rechtwinkligen Koordinatensystem die Punkte $(0, 0), (1, e_1) \dots (n, e_n)$. Ist ein Koeffizient Null, so fällt der betreffende Punkt fort. Man konstruiere nun zu den markierten Punkten das zugehörige Newtonsche Polygon, d. h. denjenigen (von unten gesehen) konvexen Polygonzug, der $(0, 0)$ als Anfangspunkt und markierte Punkte als Eckpunkte hat und so beschaffen ist, daß alle markierten Punkte auf oder über dem Polygonzug liegen. Dieser Polygonzug heiße das zum $\bmod. \mathfrak{P}$ gehörige *N. P.* von $f(x)$. Hat der Polygonzug s Seiten, so sollen ihre horizontalen und vertikalen Projektionen mit l_i, h_i ($i = 1, 2, \dots, s$) bezeichnet werden. Die Steigungen der Polygonseiten sind dann $v_i = \frac{h_i}{l_i}$ ($i = 1, 2, \dots, s$). Es ist

$$l_1 + l_2 + \dots + l_s = n, \quad v_{i+1} > v_i. \quad (3)$$

Mit diesen Bezeichnungen gilt⁴⁾

Satz 1. Es sei

$$f(x) = x^n + \alpha_1 x^{n-1} + \dots + \alpha_{n-1} x + \alpha_n = 0, \quad \alpha_n \neq 0 \quad (2)$$

eine Gleichung mit algebraischen Koeffizienten und den Wurzeln $\xi_1, \dots, \xi_n$. Ist dann $\mathfrak{P}$ ein Primideal aus dem Körper K , der $\xi_1, \dots, \xi_n$ enthält, und besitzt das zu $\mathfrak{P}$ gehörige Newtonsche Polygon s Seiten mit den horizontalen, resp. vertikalen Projektionen l_i, h_i ($i = 1, 2, \dots, s$), so sind genau l_i Wurzeln bez. $\mathfrak{P}$ von der Ordnung $v_i = \frac{h_i}{l_i}$.

Bew. Wir beweisen zunächst einen weniger präzisierten Hilfssatz. Es sei α_j der letzte Koeffizient niedrigster Ordnung $\bmod. \mathfrak{P}$ (es werde dabei $\alpha_0 = 1$ gesetzt); α_j bestimmt offenbar durch seinen Index den Endpunkt des nicht steigenden Teiles des *N. P.* Ist ein solcher nicht

⁴⁾ Der Satz dürfte zuerst von M. Bauer aufgestellt und bewiesen sein. Vgl. J. f. Math. 132 (1907), p. 21, § III.

vorhanden, wird $j = 0$. Wir wollen dann zeigen, daß $n - j$ Wurzeln durch $\mathfrak{P}$ teilbar und j Wurzeln nicht durch $\mathfrak{P}$ teilbar sind. Die Behauptung ist für Polynome 1. Grades evident, wir können daher vollständige Induktion anwenden. Der Satz möge bereits für Polynome n^{ten} Grades bewiesen sein und es sei

$$\begin{aligned} g(x) &= x^{n+1} + \beta_1 x^n + \dots + \beta_{n+1} = \\ &= (x^n + \alpha_1 x^{n-1} + \dots + \alpha_n)(x - \xi), \quad \xi \neq 0 \end{aligned} \quad (4)$$

$$\text{also } \beta_i = \alpha_i - \xi \alpha_{i-1} \quad (i = 1, 2, \dots, n), \quad \beta_{n+1} = -\alpha_n \xi. \quad (5)$$

Je nachdem ξ durch $\mathfrak{P}$ teilbar ist oder nicht, ist offenbar β_j , resp. β_{j+1} der letzte Koeffizient niedrigster Ordnung mod. $\mathfrak{P}$ in $g(x)$, woraus der Hilfssatz folgt.

Wendet man nun auf (2) die Substitution $x = y\pi^v$ an, wo π eine genau durch $\mathfrak{P}$ teilbare Zahl aus K und v eine ganze rationale Zahl bedeutet, so möge sich die Gleichung $\bar{f}(y, v) = 0$ mit dem ersten Koeffizienten 1 ergeben. $\bar{f}(y, v)$ besitzt bez. $\mathfrak{P}$ ein $N. P.$, dessen s Seiten die Steigungen $v_i - v$ und die Horizontalprojektionen l_i haben. Wir nehmen zunächst an, daß die Steigungen v_i ganze Zahlen sind, was sich später als zutreffend erweisen wird.

Das Polynom $\bar{f}(y, v_1)$ besitzt bez. $\mathfrak{P}$ lauter Koeffizienten von nicht negativer Ordnung; daher sind alle Wurzeln von (2) durch $\mathfrak{P}^{v_1}$ teilbar. Der im Hilfssatz mit j bezeichnete Index hat für $\bar{f}(y, v_1)$ den Wert l_1 , daher besitzen genau l_1 Wurzeln eine Ordnung $\leq v_1$, wo nach dem Vorhergehenden nur das Gleichheitszeichen gelten kann. Die Gleichung (2) besitzt also genau l_1 Wurzeln von der Ordnung v_1 bez. $\mathfrak{P}$.

Für die Polynome $\bar{f}(y, v_2 - 1)$ und $\bar{f}(p, v_2)$ hat der Index j des Hilfssatzes resp. die Werte l_1 und $l_1 + l_2$. Daher besitzen genau l_1 Wurzeln eine Ordnung $\leq v_2 - 1$ und genau $l_1 + l_2$ Wurzeln eine Ordnung $\leq v_2$. Die Gleichung (2) hat also genau l_2 Wurzeln von der Ordnung v_2 bez. $\mathfrak{P}$. In dieser Weise kann man weiter schließen.

Daß alle v_i ganze Zahlen sind, ergibt sich so. Es mögen sich alle v_i als Brüche mit dem Nenner a darstellen lassen. Man bilde dann den Körper $(K, \sqrt[a]{\pi}) = \bar{K}$, in dem $\bar{\mathfrak{P}} = \mathfrak{P}^{\frac{1}{a}}$ Primideal ist. Für $\bar{\mathfrak{P}}$ als Modul gehen die Zahlen v_i in av_i über, sind also ganz. Nach dem Bewiesenen haben daher genau l_i Wurzeln von (2) bez. $\bar{\mathfrak{P}}$ die Ordnung av_i . Da nun die Ordnung der Wurzeln von (2) bez. $\mathfrak{P}$ eine ganze Zahl ist, muß ihre Ordnung bez. $\bar{\mathfrak{P}}$ durch a teilbar sein, d. h. alle v_i sind ganze Zahlen.

Aus Satz 1 folgt unmittelbar⁵⁾

Satz 2 (Produktsatz). Es seien $f(x)$ und $g(x)$ zwei Polynome mit Koeffizienten aus dem algebraischen Rationalitätsbereich k und dem höchsten Koeffizienten 1; ferner sei $\mathfrak{p}$ ein Primideal aus k oder einem Oberkörper von k . Man erhält dann das Newtonsche Polygon mod. $\mathfrak{p}$ für das Produkt $f(x) \cdot g(x)$, indem man die Newtonschen Polygone mod. $\mathfrak{p}$ für $f(x)$ und $g(x)$ nach wachsender Steigung der Seiten zusammensetzt.

Bew. Es seien $\xi_1, \dots, \xi_n$ die Wurzeln von $f(x)g(x) = 0$ und $\mathfrak{P}$ ein Primideal aus $K(\xi_1, \dots, \xi_n)$ oder einem Oberkörper von K , das in $\mathfrak{p}$ genau zur u^{ten} Potenz aufgeht. Für die $N. P.$ mod. $\mathfrak{P}$ folgt unser Satz dann direkt aus Satz 1. Er gilt dann auch für die $N. P.$ mod. $\mathfrak{p}$, da diese aus den $N. P.$ mod. $\mathfrak{P}$ dadurch entstehen, daß man die Vertikalprojektionen der Seiten auf den u^{ten} Teil verkürzt.

Wir brauchen noch folgenden⁶⁾

Satz 3. Es sei

$$f(x) = x^n + \alpha_1 x^{n-1} + \dots + \alpha_n = 0 \quad (6)$$

eine Gleichung mit Koeffizienten aus dem Körper k und den Wurzeln $\xi_1, \dots, \xi_n$, die in k irreduzibel ist. Es sei ferner $\mathfrak{p}$ ein Primideal aus k und das zu $\mathfrak{p}$ gehörige Newtonsche Polygon von $f(x)$ möge s Seiten mit den horizontalen, resp. vertikalen Projektionen l_i, h_i ($i=1, 2, \dots, s$) haben. Gelten dann in $k(\xi_1)$ die Zerlegungen in Primideale

$$\mathfrak{p} = \mathfrak{p}_1^{f_1} \dots \mathfrak{p}_t^{f_t}, \quad \xi_1 = \mathfrak{p}_1^{g_1} \dots \mathfrak{p}_t^{g_t} q, \quad (q, \mathfrak{p}) = 1 \quad (7)$$

so können die Quotienten $\frac{g_i}{f_i}$ nur einen der Werte $v_i = \frac{h_i}{l_i}$

haben und es kommen auch alle v_i unter den Quotienten $\frac{g}{f}$ vor. Ist speziell $(l_i, h_i) = 1$ ($i=1, 2, \dots, s$), so gilt in $k(\xi_1)$ die Zerlegung in Primideale:

$$\mathfrak{p} = \mathfrak{p}_1^{l_1} \dots \mathfrak{p}_s^{l_s}, \quad \xi_1 = \mathfrak{p}_1^{h_1} \dots \mathfrak{p}_s^{h_s} q, \quad (q, \mathfrak{p}) = 1. \quad (8)$$

Bew. Es sei $\mathfrak{P}$ ein in $\mathfrak{p}_1$ aufgehendes Primideal aus $K(\xi_1, \dots, \xi_n)$ und $\mathfrak{p} = \mathfrak{P}^u \mathfrak{Q}$, $(\mathfrak{P}, \mathfrak{Q}) = 1$. Das zu $\mathfrak{P}$ gehörige $N. P.$ von $f(x)$ erhält man

⁵⁾ Der Satz ist für den R (1) zuerst von Dumas l. c. bewiesen, später hat ihn Ö. Ore auf algebraische Zahlkörper verallgemeinert.

⁶⁾ Vergl. M. Bauer, l. c. und Ö. Ore, Acta math. 44.

aus dem zu $\mathfrak{p}$ gehörigen, indem man alle h_i mit u multipliziert. Nach Satz 1 ist dann ξ_1 genau durch $\mathfrak{P}^{uv_l}$ teilbar, wo l einen Index der Reihe 1, 2, ... s bedeutet. Daraus folgt $v_l = \frac{g_1}{f_1}$. Daß alle v_i , z. B. v_1 , unter den Quotienten $\frac{g}{f}$ vorkommen, folgt so. Es gibt nach Satz 1 eine Wurzel, die genau durch $\mathfrak{P}^{uv}$ teilbar ist. Es ist dann auch ξ_1 genau durch $\overline{\mathfrak{P}}^{uv}$ teilbar, wo $\overline{\mathfrak{P}}$ ein konjugiertes Primideal zu $\mathfrak{P}$ bezeichnet. Geht nun $\overline{\mathfrak{P}}$ in $\mathfrak{p}_m$ auf, so folgt $v_1 = \frac{g_m}{f_m}$.

Da alle v_i verschieden sind, kann man die Numerierung der $\mathfrak{p}_i$ so annehmen, daß

$$\frac{g_i}{f_i} = v_i = \frac{h_i}{l_i} \quad (i = 1, 2, \dots, s).$$

Ist nun speziell $(l_i, h_i) = 1$ ($i = 1, 2, \dots, s$), so folgt $f_i \equiv 0 \pmod{l_i}$. Da $\sum_{i=1}^s l_i = n$ und $\sum_{i=1}^s f_i \leq n$, folgt $f_i = l_i$, $g_i = h_i$ und $t = s$. Überdies sind alle Primideale $\mathfrak{p}_i$ vom ersten Relativgrad bez. k .

§ 2.

Wir können jetzt folgenden Satz über Gleichungen ohne Affekt aufstellen.

Satz 4. Es seien $\mathfrak{p}_0, \mathfrak{p}_1, \dots, \mathfrak{p}_{n-2}$ verschiedene Primideale eines algebraischen Zahlkörpers k und π_i eine genau durch $\mathfrak{p}_i$ teilbare und zu den übrigen $\mathfrak{p}_j$ teilerfremde Zahl aus k ($i, j = 0, 1, \dots, n-2$; $i \neq j$). Ist dann

$$f(x) = x^n + \pi_0^{e_{0,1}} \dots \pi_{n-2}^{e_{n-2,1}} \alpha_1 x^{n-1} + \dots + \pi_0^{e_{0,n}} \dots \pi_{n-2}^{e_{n-2,n}} \alpha_n \quad (9)$$

$$\left. \begin{aligned} e_{i,j} &= 1 \quad (j = 1, 2, \dots, n-i) \\ e_{i,n-i+1} &= 1 + \frac{l(l+1)}{2} \quad (l = 1, 2, \dots, i) \end{aligned} \right\} \quad (i = 0, 1, \dots, n-2) \quad (10)$$

so besitzt die Gleichung $f(x) = 0$ in k keinen Affekt, wenn die Zahlen α_i relativ prim zu $\mathfrak{p}_0, \mathfrak{p}_1, \dots, \mathfrak{p}_{n-2}$ aus k gewählt werden.

Bew. Soll die Gleichung $f(x) = 0$, die die Wurzeln $\xi_1, \dots, \xi_n$ haben möge, in k keinen Affekt haben, muß der Körper k

$(\xi_1, \dots, \xi_n)$ den Grad $n!$ über k haben. Dazu ist notwendig und hinreichend, daß $f(x)$ in $k, f_1(x, \xi_1) = \frac{f(x)}{x - \xi_1}$ in $k(\xi_1), f_2(x, \xi_1, \xi_2) = \frac{f_1(x, \xi_1)}{x - \xi_2}$ in $k(\xi_1, \xi_2)$ usw. $f_{n-2}(x, \xi_1, \dots, \xi_{n-2}) = \frac{f_{n-3}(x, \xi_1, \dots, \xi_{n-3})}{x - \xi_{n-2}}$ in $k(\xi_1, \dots, \xi_{n-2})$ irreduzibel ist⁷⁾. Das Polynom $f(x)$ ist nun so konstruiert, daß das zu p_0 gehörige $N. P.$ eine einzige Seite mit der Steigung $1/n$ besitzt. Ferner besitzt das zu p_1 gehörige $N. P.$ 2 Seiten mit den Steigungen $1/n-1, 1$ und den Horizontalprojektionen $n-1, 1$. Allgemein besitzt das zu p_i gehörige $N. P.$ $i+1$ Seiten mit den Steigungen $1/n-i, 1, 2, \dots, i$ und den Horizontalprojektionen $n-i, 1, 1, \dots, 1$ ($i=0, 1, \dots, n-2$)⁸⁾. $f(x)$ ist in k irreduzibel, wie aus dem $N. P.$ für p_0 folgt, das sich nicht in 2 Teile zerlegen läßt (Eisensteinsches Kriterium). Nach Satz 3 gelten nun in $k(\xi_1)$ folgende Zerlegungen in Primideale: $p = p_{i,1} p_{i,2} \dots p_{i,i} p_{i,i+1}^2, \xi_1 = p_{i,1} p_{i,1}^2$

$$\dots p_{i,i}^i p_{i,i+1} q_i, (q_i, p_i) = 1 \quad (i=1, 2, \dots, n-2) \quad (11)$$

Speziell gilt

$$p_1 = p_{1,1} p_{1,2}^{n-1}, \xi_1 = p_{1,1} p_{1,2} q_1, (q_1, p_1) = 1 \quad (12)$$

Da $f(x) = f_1(x, \xi_1)(x - \xi_1)$ und da $f(x)$ bez. $p_{1,1}$ dasselbe $N. P.$ hat wie bez. p_1 , folgt aus dem Produktsatz, daß $f_1(x, \xi_1)$ bez. $p_{1,1}$ ein $N. P.$ mit einer einzigen Seite von der Steigung $1/n-1$ hat. Daher ist $f_1(x, \xi_1)$ in $k(\xi_1)$ irreduzibel.

Allgemein folgt, da $f(x)$ bez. $p_{i,i}$ dasselbe $N. P.$ hat wie bez. p_i und da ξ_1 genau durch $p_{i,i}$ teilbar ist, daß $f_1(x, \xi_1)$ bez. $p_{i,i}$ ein $N. P.$ von i Seiten mit den Steigungen $1/n-i, 1, 2, \dots, i-1$ und den Horizontalprojektionen $n-i, 1, 1, \dots, 1$ ($i=1, 2, \dots, n-2$) hat. Die Primideale $p_{i,i}$ ($i=1, 2, \dots, n-2$) spielen daher für $f_1(x, \xi_1)$ genau die gleiche Rolle wie die Primideale $p_0, \dots, p_{n-2}$ für $f(x)$; es tritt nur $n-1$ an Stelle von n . Man kann also weiter schließen, daß $f_2(x, \xi_1, \xi_2)$ in $k(\xi_1, \xi_2), \dots, f_{n-2}(x, \xi_1, \dots, \xi_{n-2})$ in $k(\xi_1, \xi_2, \dots, \xi_{n-2})$ irreduzibel ist, womit unser Satz bewiesen ist.

Speziell gilt in R (1) folgender

Satz 5. Sind $p_0, p_1, \dots, p_{n-2}$ verschiedene rationale Primzahlen und setzt man

⁷⁾ Auf Grund dieser Überlegung hat zuerst O. Perron Gleichungen ohne Affekt hergestellt, vergl. l. c. § 2.

⁸⁾ Die Voraussetzung über die α_i , die der Einfachheit halber so angenommen wurde, ist in dem angegebenen Umfange nicht notwendig. Wesentlich ist nur, daß die $N. P.$ bez. der p_i das angegebene Verhalten zeigen.

$$f(x) = x^n + p_0^{e_{0,1}} \dots p_{n-2}^{e_{n-2,1}} a_1^{n-1} x^{n-1} + \dots p_0^{e_{0,n}} \dots p_{n-2}^{e_{n-2,n}} a_n \quad (13)$$

$$\left. \begin{aligned} e_{i,j} &= 1 \quad (j=1, 2, \dots, n-i) \\ e_{i,n-i+l} &= 1 + \frac{l(l+1)}{2} \quad (l=1, 2, \dots, i) \quad (i=0, 1, \dots, n-2), \end{aligned} \right\} \quad (14)$$

so hat die Gleichung $f(x)=0$ in $R(1)$ keinen Affekt, wenn die rationalen Zahlen a_i zu $p_0 p_1 \dots p_{n-2}$ relativ prim sind.

§ 3.

Wir können nun ohne Schwierigkeit den in der Einleitung genannten Satz über affektfreie Gleichungen beweisen.

Satz 6. Die Gleichungen n^{ten} Grades mit Koeffizienten aus einem beliebig gegebenen reellen algebraischen Rationalitätsbereich k und dem höchsten Koeffizienten 1, die in k keinen Affekt haben, liegen im Koeffizientenraum R_n überall dicht.

Bew. Es sei

$$z^n + \gamma_1 z^{n-1} + \dots + \gamma_n = 0 \quad (15)$$

eine Gleichung n^{ten} Grades mit beliebigen reellen Koeffizienten. Es seien ferner $p_0, p_1, \dots, p_{n-2}$ verschiedene Primideale aus k und es möge π_i die gleiche Bedeutung haben wie in Satz 4. Wir bezeichnen zur Abkürzung den Koeffizienten von x^{n-i} in (9), abgesehen vom Faktor α_i , mit $\bar{\pi}_i$ und setzen $\pi = \bar{\pi}_1 \bar{\pi}_2 \dots \bar{\pi}_n$. Man kann nun zu jedem ganzzahligen t ganze rationale Zahlen a_i , die zu π relativ prim sind, so bestimmen, daß

$$\left| \frac{(\pi t + 1) \gamma_i}{\bar{\pi}_i} - a_i \right| < x, \quad (i=1, 2, \dots, n) \quad (16)$$

wo x eine von t unabhängige positive Konstante bedeutet. Es folgt dann

$$\left| \gamma_i - \frac{\bar{\pi}_i a_i}{\pi t + 1} \right| < \left| \frac{x \bar{\pi}_i}{\pi t + 1} \right|, \quad (i=1, 2, \dots, n) \quad (17)$$

wo die rechte Seite für genügend großes t beliebig klein wird. Die Gleichung

$$x^n + \frac{\bar{\pi}_1 a_1}{\pi t + 1} x^{n-1} + \dots + \frac{\bar{\pi}_n a_n}{\pi t + 1} = 0 \quad (18)$$

hat nun nach Satz 4 in k keinen Affekt und liegt der willkürlich gewählten Gleichung (15) beliebig nahe.

Es ist übrigens leicht zu sehen, daß sogar schon die rational-zahligen Gleichungen n^{ten} Grades mit dem ersten Koeffizienten 1, die in k keinen Affekt haben, im Koeffizientenraum R_n überall dicht liegen. Wählt man nämlich $p_0, p_1, \dots p_{n-2}$ als verschiedene rationale Primzahlen, die zur Diskriminante von k relativ prim sind, so ist die Gleichung (13) aus Satz 5 auch in k affektfrei.

Auch für komplexe Rationalitätsbereiche kann man einen dem Satz 6 analogen Satz aufstellen. Es sei in (15)

$$\gamma_j = \alpha_j + i\beta_j \quad (j = 1, 2, \dots n).$$

Man kann dann $\alpha_1, \alpha_2, \dots \alpha_n, \beta_1, \beta_2, \dots \beta_n$ als rechtwinklige Koordinaten in einem R_{2n} interpretieren und erhält so einen Koeffizientenraum R_{2n} . Es gilt dann

Satz 7. Die Gleichungen n^{ten} Grades mit Koeffizienten aus einem beliebigen algebraischen komplexen Rationalitätsbereich k und dem ersten Koeffizienten 1, die in k keinen Affekt haben, liegen im Koeffizientenraum R_{2n} überall dicht.

Bew. Es möge $\bar{\pi}_i$ ($i = 1, 2, \dots n$) und π die analoge Bedeutung wie im Beweis zu Satz 6 haben und es sei ω eine ganze algebraische Zahl aus k von solcher Art, daß $\omega\pi$ nicht reell ist. Man kann dann zu jedem ganzzahligen t ganze rationale Zahlen b_j, a_j so bestimmen, daß a_j zu π relativ prim ist und

$$\left| \zeta_j \right| = \left| \frac{(\pi t + 1) \gamma_j}{\bar{\pi}_j} - a_j - b_j \omega \pi \right| < \varkappa, \quad (j = 1, 2, \dots n) \quad (19)$$

wo $\varkappa$ eine von t unabhängige positive Konstante bedeutet (man wähle zuerst b_j so, daß der Koeffizient des imaginären Teiles von ζ_j zwischen $+\frac{c}{2}$ und $-\frac{c}{2}$ liegt, wo c der Koeffizient von i in $\omega\pi$ ist). Es folgt dann

$$\left| \gamma_j - \frac{\bar{\pi}_j(a_j + b_j \omega \pi)}{\pi t + 1} \right| < \left| \frac{\varkappa \bar{\pi}_j}{\pi t + 1} \right|, \quad (20)$$

wo die rechte Seite für beliebig großes t beliebig klein wird. Die Gleichung

$$x^n + \frac{\bar{\pi}_1(a_1 + b_1 \omega \pi)}{\pi t + 1} x^{n-1} + \dots + \frac{\bar{\pi}_n(a_n + b_n \omega \pi)}{\pi t + 1} = 0 \quad (21)$$

hat dann nach Satz 4 in k keinen Affekt und liegt der willkürlich gewählten Gleichung (15) beliebig nahe, womit unsere Behauptung bewiesen ist.

(Eingegangen: 12. VI. 1928.)

Beiträge zur mehrdimensionalen Differential- geometrie ¹⁾.

Von C. Burstin in Wien.

Die Hypertorse und ihr verwandte Hyperflächen des
euklidischen R_n .

Wir wollen hier die Hyperflächen F_l :

$$(1) \quad x_i = x_i(y_1, y_2, \dots, y_l), \quad i = 1, 2, \dots, n$$

untersuchen, deren J_2 -Raum eindimensional ist und deren Riemannscher Krümmungstensor verschwindet. Als Koordinaten des euklidischen R_n wählen wir rechtwinklige kartesische. Wenn ξ_i der den J_2 -Raum aufspannende Einheitsvektor ist, so gilt für die zweiten Ableitungen der x_i nach den y_t die Darstellung:

$$(2) \quad \frac{\partial^2 x_i}{\partial y_p \partial y_q} = \left\{ \begin{matrix} p & q \\ t \end{matrix} \right\} \frac{\partial x_i}{\partial y_t} + B_{pq} \xi_i; \quad p, q, t = 1, 2, \dots, l; \quad i = 1, 2, \dots, n.$$

Dabei sind $\left\{ \begin{matrix} p & q \\ l \end{matrix} \right\}$ die Christoffelklammern der Maßform

$$(3) \quad \gamma_{pq} = \frac{\partial x_i}{\partial y_p} \frac{\partial x_i}{\partial y_q} \text{ der } F_l.$$

Infolge unserer zweiten Voraussetzung des Verschwindens des Krümmungstensors gilt für die Form B_{pq}

$$(4) \quad B_{pq} B_{rs} - B_{ps} B_{rq} = 0, \text{ woraus sofort } (4') \quad B_{pq} = a_p \cdot a_q \text{ folgt } ^2).$$

Die Pfaffsche Gleichung

$$(5) \quad a_p dy_p = 0 \text{ resp. } B_{pq} dy_p = 0 \text{ besitzt einen Multiplikator, da}$$

$$(6) \quad \left(\frac{\partial B_{pq}}{\partial y_r} - \frac{\partial B_{pr}}{\partial y_q} \right) B_{pt} + \left(\frac{\partial B_{pr}}{\partial y_t} - \frac{\partial B_{pt}}{\partial y_r} \right) B_{pq} + \left(\frac{\partial B_{pt}}{\partial y_q} - \frac{\partial B_{pq}}{\partial y_t} \right) B_{pr} = 0$$

¹⁾ Diese Abhandlung bildet den zweiten von C. Burstin ausgearbeiteten Teil einer mit W. Mayer gemeinsam verfaßten Arbeit, deren erster Teil, von W. Mayer ausgearbeitet, unter dem Titel „Über das vollständige Formensystem der F_e im R_n “ in den Monatsheften für Mathematik und Physik, XXXV. Band, veröffentlicht wurde.

²⁾ Kanonische Darstellung des Tensors B_{pq} vom Range 1.

aus den Integrabilitätsbedingungen von (2) und aus (4') folgt. Der Koeffizient ξ_i der Integrabilitätsbedingungen von (2) lautet

$$(7) \quad \left\{ \begin{smallmatrix} p & q \\ h \end{smallmatrix} \right\} B_{hr} - \left\{ \begin{smallmatrix} p & r \\ h \end{smallmatrix} \right\} B_{hq} + \left(\frac{\partial B_{pq}}{\partial y_r} - \frac{\partial B_{pr}}{\partial y_q} \right) = 0, \quad \text{was dann (6) als}$$

Folge hat.

Es existiert somit ein Integral $\varphi(y_1 y_2 \dots y_e) = \text{konst.}$ der Gleichung

$$a_p dy_p = 0.$$

Führen wir neue Parameter $z_1, z_2, \dots, z_e$ auf der F_e ein, daß $z_1 = \varphi(y_1 y_2 \dots y_e)$ dieses Integral ist, wobei wir uns die Fixierung der $z_2, \dots, z_e$ noch vorbehalten, und behalten wir für diese neuen Parameter die Bezeichnungen $a_p, B_{pq}, \left\{ \begin{smallmatrix} p & q \\ t \end{smallmatrix} \right\}, \gamma_{pq}$ bei, so folgt aus der Tatsache, daß $a_p \cdot dz_p = 0$ jetzt $z_1 = \text{konst.}$ zum Integral hat:

$$(8) \quad a_i = 0, \quad i \neq 1 \quad \text{und} \quad B_{ik} = 0, \quad i \quad \text{oder} \quad k \neq 1.$$

Statt (2) gilt nun:

$$(9) \quad \frac{\partial^2 x_i}{\partial z_p \partial z_q} = \left\{ \begin{smallmatrix} p & q \\ t \end{smallmatrix} \right\} \frac{\partial x_i}{\partial z_t} + B_{pq} \xi_i, \quad \text{wo jetzt nur } B_{11} \neq 0 \text{ ist.}$$

Betrachten wir in unserer F_l

(10) $z_i = z_i(z_1 z_2 \dots z_l)$, die $E_{l-1} z_1 = \text{konst.}$, so ist diese eine Hyperebene E_{e-1} .

Für $p \neq 1, q \neq 1$ folgt aus (7)

$$(11) \quad \left\{ \begin{smallmatrix} p & q \\ 1 \end{smallmatrix} \right\} = 0, \quad p \neq 1, \quad q \neq 1, \quad \text{also gilt für unsere } F_{l-1}$$

$$(12) \quad \frac{\partial^2 x_i}{\partial z_p \partial z_q} = \left\{ \begin{smallmatrix} p & q \\ t \end{smallmatrix} \right\} \frac{\partial x_i}{\partial z_t}; \quad p, q, t = 2, 3 \dots l; \quad i = 1, 2, \dots n.$$

Der erste Normalraum dieser F_{l-1} verschwindet, sie ist also eine $E_{l-1}, q. e. d.$

Wir können demnach durch Transformation der $z_2, z_3 \dots z_e$ in dieser E_{l-1} kartesische Koordinaten einführen mit dem Erfolg, daß alle $\left\{ \begin{smallmatrix} p & q \\ t \end{smallmatrix} \right\} = 0; p, q, t = 2, \dots, l$ verschwinden und daß dann

$$(12') \quad \frac{\partial^2 x_i}{\partial z_p \partial z_q} = 0; \quad p, q = 2, \dots, l \text{ lautet.}$$

Dies hat dann für unsere F_l die Darstellung

$$(13) \quad x_i = A_i(z_1) + C_{ik}(z_1)z_k; \quad k = 2, 3, \dots, l$$

als Folge, wo die Matrix $\|C_{ik}\|$ vom Range $l-1$ sein muß (sonst hätten wir keine F_l vor uns), und wo wir die C_{ik} ; $k = 2, \dots, l$; noch als orthogonales $l-1$ -Bein voraussetzen dürfen³⁾. Wir können durch eine orthogonale Transformation (alle Indizes von 2 bis l)

$$(14) \quad z_k = a_{kr} \bar{z}_r \quad (14') \quad a_{kr} a_{ks} = \delta_{rs}$$

$$a_{kr} = a_{kr}(z_1) \text{ stets } \bar{C}_{ik} \bar{C}_{ij}' = 0 \quad i = 1, 2, \dots, n, \text{ erreichen.}$$

In der Tat gilt ja

$$(15) \quad \bar{C}_{ir} = C_{ik} a_{kr}, \quad \bar{C}_{is}' = C_{ij}' a_{js} + C_{ij} a_{js}', \text{ also} \\ \bar{C}_{ir} \bar{C}_{is}' = a_{jr} a_{js}' + C_{ik} C_{ij}' a_{kr} a_{js} = 0,$$

wo in $C_{ik} C_{ij}' = \delta_{kj}$ die Beineigenschaft der C_{ik} verwendet wurde. Multipliziert man die letzte Gleichung mit a_{hr} , so erhält man

$$(15') \quad a_{hs}' + C_{ih} C_{ij}' a_{js} = 0$$

als System Differentialgleichungen zur Bestimmung der a_{js} . Ein Integral des Systems (15') ist $a_{hr} a_{hs} = \text{konst. (rs)}$, somit können wir ein (14') genügendes Lösungssystem a_{hs} von (15') finden. Aus (15) folgt, daß die $\bar{C}_{ik}$ wieder ein orthogonales $l-1$ -Bein bilden.

Wir denken uns, daß $z_2, \dots, z_l$ selbst dieses gestrichene System ($\bar{z}_h$; $h = 2, \dots, l$) ist, dann gilt also neben (13):

$$(16) \quad C_{ih} \bar{C}_{ij}' = 0; \quad i = 1, 2, \dots, n; \quad k, j = 2, \dots, l.$$

Führen wir jetzt statt $A_i: A_i = A_i + C_{ik} w_k(z_1)$ in (13) ein, so hat dies wieder

$$(17) \quad x_i = A_i(z_1) + C_{ik}(z_1)\{z_k - w_k(z_1)\} = \bar{A}_i(z_1) + C_{ik} z_k; \quad k = 2, \dots, l$$

³⁾ In der Tat: Aus $\frac{\partial x_i}{\partial z_k} = C_{ik}$ und $\frac{\partial x_i}{\partial z_k} \frac{\partial x_i}{\partial z_j} = C_{ik} C_{ij} = \gamma_{kj}(k, j = 2, \dots, l)$

folgt, daß die C_{ik} ein solches orthogonales $l-1$ -Bein bilden, wenn wir in der E_{l-1} rechtwinklige kartesische Koordinaten einführen.

als Folge, und wir können noch die $w_k(z_1)$ so bestimmen, daß

$$A'_i C_{ij} = [A'_i + C'_{ik} w_k + C_{ik} w'_k] C_{ij} = A'_i C_{ij} + w'_j = 0 \text{ ist.}$$

Dies möge bereits für unsere Koordinaten $z_1 z_2 \dots z_l$ zutreffen, so daß jetzt neben (16) noch

$$(18) \quad A'_i C_{ij} = 0 \text{ zutrifft.}$$

Ferner notieren wir

$$(19) \quad C_{ik} C_{ij} = \delta_{ij}; \quad k, j = 2, \dots, l.$$

Jetzt haben wir die $z_2, \dots, z_l$, deren Fixierung wir uns vorbehalten, bestimmt.

Der Tangentialraum der F_l (13) wird durch die l Vektoren

$$(20) \quad \begin{cases} \frac{\partial x_l}{\partial z_1} = A'_l + C'_{lk}(z_1) z_k \\ \frac{\partial x_i}{\partial z_p} = C_{ip}(z_1) \end{cases}; \quad k, p = 2, \dots, l$$

aufgespannt. Ferner gilt $\frac{\partial^2 x_l}{\partial z_1 \partial z_p} = C'_{ip}(z_1); B_{1p} = 0$ hat also

$$(21) \quad C'_{ip} = \alpha_p [A'_i + C'_{ik} z_k] + \alpha_{pq} C_{iq}(z_1); \quad (p, q = 2, \dots, l)$$

als Folge.

Bezeichnen wir $\alpha_p|_{z_2=z_3=\dots=z_l=0} = \varphi_p$, $\alpha_{pq}|_{z_2=\dots=z_l=0} = \varphi_{pq}$, so führt $z_2=z_3=\dots=z_l=0$ in (21) zu:

$$(22) \quad C'_{ip} = \varphi_p A'_i + \varphi_{pq} C_{iq},$$

wo Multiplikation mit C_{ir} (16), (18) und (19):

$$(23) \quad \varphi_{pr} = 0; \quad p, r = 2, \dots, l \text{ ergibt.}$$

Es gilt somit die wichtige Formel:

$$(24) \quad C'_{ip} = \varphi_p A'_i {}^4)$$

Hier wurde $\alpha_p(z_1 z_2 \dots z_l)|_{z_2=z_3=\dots=z_l=0} = \varphi_p(z_1) =$ endlicher Wert vorausgesetzt. Die Diskussion der Nichtendlichkeit bringen wir am Schlusse.

⁴⁾ Gilt (24), so ist $\frac{\partial^2 x_i}{\partial z_1 \partial z_p} = C'_{ip} = \varphi_p A'_i = \varphi_p \pi \frac{\partial x_i}{\partial z_1}; \frac{\partial^2 x_i}{\partial z_p \partial z_q} = 0; \quad p, q \neq 1$, d. h. der J_2 Raum ist höchstens eindimensional. Da weiter nur $B_{11} \neq 0$ ist, so verschwindet die Riemannsche Krümmung.

Wir führen in unserer F_l neue Parameter $\bar{z}_1, \bar{z}_2, \dots, \bar{z}_l$ ein

$$(25) \quad z_1 = \bar{z}_1, \quad z_k = \bar{z}_k + c_k(z_1),$$

wobei wir uns die Bestimmung der $c_k(z_1)$ vorbehalten.

Die F_l lautet in den neuen Parametern

$$(26) \quad x_i = \bar{A}_i(\bar{z}_1) + \bar{O}_{ik}(\bar{z}_1)\bar{z}_k; \quad k = 2, \dots, l,$$

wobei

$$(26') \quad \bar{C}_{ik}(\bar{z}_1) = C_{ik}(z_1) \text{ und } \bar{A}_i(\bar{z}_1) = A_i + C_{ik}c_k(z_1) \text{ gilt.}$$

Aus (26') und (24) folgt:

$$\bar{A}_i' = A_i' + c_k C_{ik}' + c_k' C_{ik} = A_i'(1 + \varphi_k c_k) + c_k' C_{ik}.$$

Setzen wir als erste Bestimmungsgleichung der c_k

$$(27) \quad 1 + \varphi_k c_k = 0, \text{ so hat dies (27')} \quad \bar{A}_i' = c_k' C_{ik} \text{ als Folge.}$$

Aus (27') und (24) folgt:

$$\bar{A}_i'' = c_k'' C_{ik} + c_k' C_{ik}' = c_k'' C_{ik} + \varphi_k c_k' A_i'.$$

Setzen wir als zweite Bestimmungsgleichung der c_k

$$(28) \quad \varphi_k c_k' = -\varphi_k' c_k = 0, \text{ so hat dies (28')} \quad \bar{A}_i'' = c_k'' C_{ik}$$

als Folge.

Ebenso folgt aus

$$(29) \quad \varphi_k c_k'' = -\varphi_k' c_k' = \varphi_k'' c_k = 0, \quad (29') \quad \bar{A}_i''' = c_k''' C_{ik} \text{ usw.}$$

bis zu

$$(30) \quad \varphi_k c_k^{(l-2)} = -\varphi_k' c_k^{(l-3)} = \dots = \pm c_k \varphi_k^{(l-2)} = 0,$$

das (30') $\bar{A}_i^{(l-1)} = c_k^{(l-1)} C_{ik}$ als Folge hat. Für die c_k haben wir damit $l-1$ Gleichungen

$$(31) \quad 1 + c_k \varphi_k = 0, c_k \varphi_k' = 0, c_k \varphi_k'' = 0, \dots; c_k \varphi_k^{(l-2)} = 0; k = 2, 3, \dots, l$$

aus denen die c_k berechnet werden können, sobald die Determinante

$$(32) \quad \Delta' = \begin{vmatrix} \varphi_2, & \varphi_3 \dots & \varphi_l \\ \varphi_2', & \varphi_3' \dots & \varphi_l' \\ \vdots & \vdots & \vdots \\ \varphi_2^{(l-2)}, & \varphi_3^{(l-2)} \dots & \varphi_l^{(l-2)} \end{vmatrix} \neq 0 \text{ ist.}$$

Mit Δ_1 bezeichnen wir die Determinante

$$(33) \quad \Delta_1 = \begin{vmatrix} c_2', & c_3' \dots & c_l' \\ c_2'', & c_3'' \dots & c_l'' \\ \vdots & \vdots & \vdots \\ c_2^{(l-1)}, & c_3^{(l-1)} \dots & c_l^{(l-1)} \end{vmatrix}$$

dann folgt nach einfacher Rechnung aus (28), (29), (30):

$$(34) \quad \Delta \Delta_1 = \pm (\varphi_k^{(l-2)} c_k')^{l-1}$$

I. Fall: $\varphi_k^{(l-2)} c_k' \neq 0$, dann ist Δ und Δ_1 von Null verschieden und wir können aus (31) die c_k und hierauf aus (27'), (28'), (29') und (30') die C_{ik} linear durch die $\overline{A}_i'$, $\overline{A}_i''$, ..., $\overline{A}_i^{(l-1)}$ berechnen:

$$(35) \quad C_{ik} = p_{kr} \overline{A}_i^{(r)}; \quad r = 1, 2, \dots, l-1, \quad k = 2, \dots, l.$$

Dies in (26) eingeführt ergibt

$$(36) \quad x_i = \overline{A}_i + p_{kr} \overline{A}_i^{(r)} \overline{z}_k$$

und nach einer letzten Parametertransformation

$$(37) \quad u_1 = z_1, \quad u_{r+1} = p_{kr} \overline{z}_k; \quad k = 2, 3, \dots, l; \quad r = 1, 2, \dots, l-1.$$

Die F_l (37) ist die Gesamtheit der Punkte aller $l-1$ dimensionalen Schmiegungs- E_{l-1} der Raumkurve $x_i = \overline{A}_i(u_1)$ des R_n . (Eine $l-1$ dimensionale Schmiegungebene in einem Punkte einer C_1 ist die E_{l-1} , die durch Tangente, erste, zweite ... bis $l-2$ te Normale dieser C_1 in diesem Punkte aufgespannt wird.) Eine solche F_l nennen wir eine Hypertorse.

II. Fall. $C_k^{(l-2)} c_k' = 0$

Unterfall a) $\Delta \neq 0$, $\Delta_1 = 0$. Differentiation der Glieder des Systems (31) ergibt dann

$$(38) \quad c_k' \varphi_k = 0, \quad c_k' \varphi_k' = 0, \dots, c_k' \varphi_k^{(l-2)} = 0,$$

aus welchem System wegen $\Delta_k \neq 0$ $c'_k = 0$ folgt. Dann ist (27') auch $\bar{A}'_i = 0$, also $\bar{A}_i = M_i = \text{konst.}$ (26') lautet dann

$$(39) \quad A_i = -C_{ik} c_k + M_i; \quad c_k, \quad M_i (k = 2, \dots, l; \quad i = 1, 2, \dots, n)$$

konstant.

Daraus und aus (24) resultiert:

$$(40) \quad C'_{ik} = -\varphi_k c_j C'_{ij}, \quad c_j \text{ konst.}; \quad k = 2, 3, \dots, l; \quad j = 2, 3, \dots, l.$$

Wir führen in (26) neue Flächenparameter ein:

$$(41) \quad \bar{z}_1 = v_1, \quad \bar{z}_2 = c_2 v_2, \quad \bar{z}_j = v_j + c_j v_2; \quad j = 3, 4, \dots, l,$$

und erhalten:

$$(42) \quad x_i - M_i = C_{i2} c_2 v_2 + \sum_{j=3}^l C_{ij} (v_j + c_j v_2) = \bar{v}_2 (\tilde{C}_{i2} + \bar{v}_j C_{ij}), \text{ wo}$$

$$(43) \quad \bar{v}_2 = v_2, \quad \bar{v}_j = \frac{v_j}{v_2} \text{ und } \tilde{C}_{i2} = C_{ik} c_k; \quad k = 2, 3, \dots, l.$$

bedeutet.

Dabei tritt an Stelle von (40)

$$(44) \quad C'_{ik} = -\varphi_k \tilde{C}'_{i2}; \quad k = 2, 3, \dots, l.$$

Für die F_{l-1}

$$(45) \quad \tilde{x}_i = \tilde{C}_{i2} + C_{ij} \bar{v}_j, \quad i = 1, 2, \dots, n; \quad j = 3, 4, \dots, l$$

sind (44) die (24) entsprechenden Beziehungen der F_l (13). Für unsere F_l selbst gilt (42)

$$(46) \quad x_i - M_i = \bar{v}_2 \tilde{x}_i; \quad i = 1, 2, \dots, n$$

d. h. unsere F_l ist ein l dimensionaler Hyperkegel mit der Spitze M_i und der $l-1$ dimensionalen Leithyperfläche $\tilde{x}_i$.

Dabei ist diese Leithyperfläche $\tilde{x}_i$ eine F_{l-1} , derselben charakteristischen Eigenschaft [wegen (44)] wie die F_l selbst. (Riemannsche Krümmung Null, J_2 Raum eindimensional.)

Unterfall $b)$ $\Delta = 0$, d. h., daß zwischen den $l-1$ Funktionen $\varphi_2, \varphi_3, \dots, \varphi_l$ eine lineare homogene Relation mit konstanten Koeffizienten besteht, in der nicht alle Koeffizienten verschwinden.

$$(47) \quad \varphi_k a_k = 0; \quad k = 2, \dots, l \quad (a_k \text{ konst.}).$$

Aus (24) folgt dann

$$(48) \quad C'_{ip} a_p = 0, \text{ also } (48') \quad C_{ip} a_p = D_i \quad (a_p, D_i \text{ konstante}).$$

Sei, was durch Umnumerierung stets erreichbar ist, $a_l \neq 0$, so schreiben wir (48'):

$$(49) \quad C_{il} = \bar{D}_i - \sum_{j=2}^{l-1} C_{ij} b_j, \text{ wobei } b_j = \frac{a_j}{a_l}, \quad \bar{D}_i = \frac{D_i}{a_l} \text{ ist.}$$

Setzen wir diese Werte in (13) ein, so erhalten wir

$$(50) \quad x_i = A_i + \sum_{j=2}^{l-1} C_{ij} (z_j - b_j z_l) + \bar{D}_i z_l.$$

Eine Parametertransformation:

$$\bar{z}_j = z_j - b_j z_l, \quad \bar{z}_1 = z_1, \quad \bar{z}_l = z_l; \quad j = 2, 3, \dots, l-1$$

hat endlich

$$(51) \quad x_i = A_i + \sum_{j=2}^{l-1} C_{ij} \bar{z}_j + \bar{D}_i \bar{z}_l$$

zum Resultat.

Setzen wir

$$(52) \quad \tilde{x}_i = A_i + C_{ij} \bar{z}_j; \quad j = 2, 3, \dots, l-1,$$

so ist dies wegen (24) eine Hyperfläche F_{l-1} , für die wegen (24) genau dieselben Voraussetzungen gelten wie für unsere F_l .

Für diese selbst gilt dann

$$(53) \quad x_i = \tilde{x}_i + \bar{D}_i \bar{z}_l; \quad D_i \text{ konst.},$$

d. h. unsere F_e ist ein l dimensionaler Hyperzylinder mit der $l-1$ dimensionalen Leitfläche $\tilde{x}_i$.

Wir haben nur mehr die Voraussetzung der Nichtendlichkeit der

$$\alpha_p(z_1 z_2 \dots z_l)_{z_2=z_3=\dots z_l=0} = \varphi_p(z_1),$$

deren Negation uns von (21) zu (24) führte, zu diskutieren:

Es sei also $\alpha_p(z_1 z_2 \dots z_l)_{z_2=\dots z_l=0}$ nicht endlich.

Multiplizieren wir (21) mit C_{ir} , so folgt $\alpha_{pr} = 0$ und aus

$$C'_{ip} = \alpha_p (z_1 z_2 \dots z_l) [A'_i + C'_{ik} z_k] \text{ folgt (für } z_2 = z_3 = \dots = z_l = 0) \\ A'_i = 0, \text{ d. h. } A_i = M_i \text{ (konstant).}$$

Dies hat

$$(54) \quad C'_{ip} = \alpha_p z_k C'_{ik}; \quad p, k = 2, 3, \dots, l$$

als Folge.

Sind alle $C'_{ik} = 0$, d. h. die $C_{ik} = \text{konst.}$, so ist nach (13) unsere F_l eine l dimensionale Hyper- E_l , ihr J_2 Raum verschwindet, d. h. aber, daß infolge unserer Voraussetzung dieser Fall nicht eintritt.

Es muß also ein C'_{ik} z. B. $C'_{i2} \neq 0$ sein, $i = 1, 2, \dots, n$ (d. h. nicht alle Komponenten von C'_{i2} dürfen verschwinden). Setzen wir dann in (54) $z_2 = 1, z_3 = z_4 = \dots = z_l = 0$, so gilt

$$(55) \quad C'_{ip} = \alpha_p (z_1, 1, 0, \dots, 0) C'_{i2}; \quad p = 2, 3, \dots, l.$$

Schreiben wir jetzt (13):

$$(56) \quad x_i = M_i + z_2 \left[C_{i2}(z_1) + \sum_{k=3}^l C_{ik}(z_1) \bar{z}_k \right]; \quad \bar{z}_k = \frac{z_k}{z_2},$$

so ist die F_{l-1}

$$(57) \quad \tilde{x}_i = C_{i2}(z_1) + \sum_{k=3}^l C_{ik}(z_1) \bar{z}_k,$$

infolge (55) eine F_{l-1} derselben charakteristischen Eigenschaft, wie die vorgegebene F_l , und diese selbst lautet dann:

$$(58) \quad x_i - M_i = z_2 \tilde{x}_i$$

d. h. unsere F_l ist ein Hyperkegel mit der Spitze M_i und der $l-1$ dimensionalen Leithyperfläche $\tilde{x}_i$ (analog Fall IIa).

Die untersuchte F_l ist daher entweder eine Hypertorse H_l oder ein Hyperkegel K_l oder ein Hyperzylinder C_l , deren Leithyper F_{l-1} entweder Hypertorse H_{l-1} oder Hyperkegel K_{l-1} oder Hyperzylinder C_{l-1} ist usw.

Im Falle der Hypertorse ist das Problem erschöpft, im anderen Falle (Zylinder oder Kegel) um eine Dimension reduziert, ohne aber eine weitere Untersuchung (außer einer Klassifikation all dieser F_l) mehr nötig zu machen. Unsere F_l entsteht allgemein aus einer H_r , $r = l$ (r -dimensionalen Hypertorse) durch Dimensionsvergrößerung mittels fortgesetzter Kegel- resp. Zylinderbildung, bis die Dimension l erreicht ist.

Was die Umkehrung unseres Resultates betrifft, beweisen wir, daß eine Hypertorse H_r einen eindimensionalen J_2 Raum und eine verschwindende Riemannsche Krümmung besitzt. Dies ist gezeigt, wenn für die H_r

$$(59) \quad x_i = A_i(z_1) + A_i^{(h)}(z_1)z_{h+1}; \quad h = 1, 2, \dots, r-1$$

bewiesen wird, daß alle zweiten Ableitungen $\frac{\partial^2 x_i}{\partial z_p \partial z_q}$ mit Ausnahme von $\frac{\partial^2 x_i}{\partial z_1^2}$ bereits im Tangentialraum J_1 der H_r liegen.

Dann ist der J_2 eindimensional und wegen $B_{pq} = 0$ mit Ausnahme von $B_{11} \neq 0$ verschwindet die Riemannsche Krümmung der H_r . Der Tangentialraum J_1 der H_r (59) wird durch die r -Vektoren aufgespannt

$$(60) \quad J_1 \left\{ \begin{array}{l} \frac{\partial x_i}{\partial z_1} = A'_i + A_i^{(h+1)} z_{h+1}; \quad h = 1, 2, \dots, r-1 \\ \frac{\partial x_i}{\partial z_{h+1}} = A_i^{(h)} \end{array} \right.$$

Demnach ist

$$\begin{aligned} \frac{\partial^2 x_i}{\partial z_{h+1} \partial z_{k+1}} &= 0; \quad h, k = 1, \dots, r-1; \quad \frac{\partial^2 x_i}{\partial z_1 \partial z_{h+1}} = A_i^{(h+1)} = \frac{\partial x_i}{\partial z_{h+1}}; \\ h &= 1, 2, \dots, r-2; \quad \frac{\partial^2 x_i}{\partial z_1 \partial z_r} = A_i^{(r)} = \frac{1}{z_r} \times \\ &\times \left[\frac{\partial x_i}{\partial z_1} - \frac{\partial x_i}{\partial z_2} - z_2 \frac{\partial x_i}{\partial z_3} - \dots - z_{r-1} \frac{\partial x_i}{\partial z_r} \right] \end{aligned}$$

und unsere Behauptung ist nachgewiesen.

Weiter zeigen wir: Liegen für eine t -dimensionale Hyperfläche F_t

$$\tilde{x}_i = \tilde{x}_i(z_1, z_2, \dots, z_t)$$

bereits alle Vektoren $\frac{\partial^2 \tilde{x}_i}{\partial z_p \partial z_q}$, p oder $q \neq 1$ im Tangentialraum

$$\tilde{J}_1 \left\{ \frac{\partial \tilde{x}_i}{\partial z_p}; \quad p = 1, 2, \dots, t \right\},$$

so gilt für den

$t+1$ dimens. Hyperkegel $x_i - M_i = \tilde{x}_i z_{t+1}$
 resp. für den $t+1$ dimens. Hyperzylinder: $x_i = \tilde{x}_i + M_i z_{t+1}$ $\left. \vphantom{\begin{array}{l} x_i - M_i = \tilde{x}_i z_{t+1} \\ x_i = \tilde{x}_i + M_i z_{t+1} \end{array}} \right\} M_i = \text{konst.}$
 das nämliche.

In der Tat wird J_1 der Tangentialraum des Hyperkegels, durch die Vektoren des $\tilde{J}_1$ und den Vektor $\tilde{x}_i$ aufgespannt:

$$(61) \quad \frac{\partial x_i}{\partial z_p} = z_{t+1} \frac{\partial \tilde{x}_i}{\partial z_p}; \quad p = 1, 2, \dots, t; \quad \frac{\partial x_i}{\partial z_{t+1}} = \tilde{x}_i.$$

Weiter gilt

$$\frac{\partial^2 x_i}{\partial z_p \partial z_q} = z_{t+1} \frac{\partial^2 \tilde{x}_i}{\partial z_p \partial z_q} = \text{Vektor des } \tilde{J}_1, \text{ für } p \text{ oder } q \neq 1;$$

$$p \text{ und } q \neq t+1. \quad \frac{\partial^2 x_i}{\partial z_{t+1} \partial z_p} = \frac{\partial \tilde{x}_i}{\partial z_p}, \quad \frac{\partial^2 x_i}{(\partial z_{t+1})^2} = 0,$$

was unsere Behauptung für den Hyperkegel beweist.

Für den Hyperzylinder gilt:

$$(62) \quad \frac{\partial x_i}{\partial z_p} = \frac{\partial \tilde{x}_i}{\partial z_p}; \quad p = 1, 2, \dots, t, \quad \frac{\partial x_i}{\partial z_{t+1}} = M_i$$

$$\frac{\partial^2 x_i}{\partial z_p \partial z_q} = \frac{\partial^2 \tilde{x}_i}{\partial z_p \partial z_q} = \text{Vektor des } \tilde{J}_1 \text{ für } p \text{ oder } q \neq 1; \quad p \text{ und } q \neq t+1$$

$$\frac{\partial^2 x_i}{\partial z_{t+1} \partial z_p} = \frac{\partial^2 x_i}{(\partial z_{t+1})^2} = 0, \text{ was unsere Behauptung für den Hyper-} \\ \text{zylinder beweist.}$$

Demnach erhalten wir durch diese Zylinder- resp. Kegelbildung aus der gegebenen H_r stets Flächen, für die alle $\frac{\partial^2 x_i}{\partial z_p \partial z_q}$ im J_1 Raume liegen mit Ausnahme von $\frac{\partial^2 x_i}{\partial z_1^2}$, für den dies nicht zutrifft. Alle diese Flächen haben aber dann nach dem eben gezeigten einen eindimens. J_2 -Raum und verschwindende Riemannsche Krümmung.

II. Ausgezeichnete Koordinatensysteme im S_n und einige Folgerungen.

Wie im euklidischen Raume R_n , so gibt es auch in Räumen konstanter Krümmungen S_n jene ausgezeichneten Koordinatensysteme, in denen die Geraden durch lineare Gleichungen in den Koordinaten x_i gegeben sind.

Unseren Betrachtungen sei ein solches Koordinatensystem zugrunde gelegt:

$$(63) \quad x_i = a_i t + b_i^5); \quad i = 1, 2, \dots, n$$

mit den Konstanten a_i, b_i stellt dann eine Gerade dar, genügt somit dem Systeme der Differentialgleichungen:

$$(63') \quad \frac{d^2 x_i}{ds^2} + \left\{ \begin{matrix} pq \\ i \end{matrix} \right\} \frac{dx_p}{ds} \frac{dx_q}{ds} = 0,$$

wobei s die Bogenlänge $ds^2 = g_{ik} dx_i dx_k$ bedeutet, und g_{ik} der Maßtensor des S_n ist, für den (63) die Geraden sind und $\left\{ \begin{matrix} pq \\ i \end{matrix} \right\}$ die Christoffelklammern für den Tensor g_{ik} sind.

Gleichzeitig gilt längs der Geraden (63):

$$(64) \quad g_{ik} a_i a_k \left(\frac{dt}{ds} \right)^2 = 1, \quad (64) \text{ nach } s \text{ differenziert ergibt}$$

$$(64') \quad \frac{\partial g_{ik}}{\partial x_j} a_i a_k a_j \left(\frac{dt}{ds} \right)^2 + 2 g_{ik} a_i a_k \frac{d^2 t}{ds^2} = 0.$$

Aus (64) und (64') folgt:

$$(65) \quad a_i \frac{d^2 t}{ds^2} + \left\{ \begin{matrix} pq \\ i \end{matrix} \right\} a_p a_q \left(\frac{dt}{ds} \right)^2 = 0.$$

Setzen wir in (65) $a_i = 0$ und $a_h \neq 0$, wenn $h \neq i$ ist, so erhalten wir:

$$(66) \quad \left\{ \begin{matrix} pq \\ i \end{matrix} \right\} = 0, \quad \text{für } p \neq i, q \neq i.$$

Aus (64'), (65), (66) bekommt man:

$$(67) \quad 2 g_{rs} a_r a_s \left[2 \sum_{p \neq i} \left\{ \begin{matrix} pi \\ i \end{matrix} \right\} a_p + \left\{ \begin{matrix} ii \\ i \end{matrix} \right\} a_i \right] = \frac{\partial g_{ii}}{\partial x_j} a_i a_k a_j.$$

Setzen wir in (67) $a_i \neq 0, a_p = 0$ für $i \neq p$, so bekommen wir:

$$(68) \quad \frac{\partial g_{ii}}{\partial x_i} = 2 g_{ii} \left\{ \begin{matrix} ii \\ i \end{matrix} \right\}, \quad \text{d. h. } \left\{ \begin{matrix} ii \\ i \end{matrix} \right\} = \frac{1}{2} \frac{\partial \ln g_{ii}}{\partial x_i} = \Phi_i.$$

Substituieren wir in (67) $a_p \neq 0, a_q = 0$, für $p \neq q$, wobei auch $i \neq p$ ist, so erhalten wir:

$$(69) \quad 4 g_{pp} \left\{ \begin{matrix} pi \\ i \end{matrix} \right\} = \frac{\partial g_{pp}}{\partial x_p}, \quad \text{d. h. } \left\{ \begin{matrix} pi \\ i \end{matrix} \right\} = \frac{1}{4} \frac{\partial \ln g_{pp}}{\partial x_p} = \frac{1}{2} \Phi_p.$$

⁵⁾ Allgemein stellt dann jede Gleichung $x_i = A_{ik} t_k + B_i$ ($k = 1, 2, \dots, l$), A_{ik}, B_i konstant eine l -dimensionale Hyperebene E_l im S_n dar.

Bezeichnen wir jetzt mit $\vartheta \lambda^i$:

$$(70) \quad \vartheta \lambda^i = d\lambda^i + \left\{ \begin{matrix} p & q \\ i \end{matrix} \right\} \lambda^p dx_q$$

die Riccische Ableitung des Vektors λ^i , dann ist für unseren Tensor g_{ik} und für das Koordinatensystem (61)

$$(70') \quad \vartheta \lambda^i = d\lambda^i + (\Phi_p dx_p) \lambda^i + (\Phi_p \lambda^p) dx_i.$$

Im S_n , dem das Koordinatensystem (61) zugrunde liegt, sei eine Kurve C ,

$$(71) \quad x_i = x_i(t) \text{ gegeben.}$$

Die Tangente J_1 von (71) ist aufgespannt durch den Vektor dx_i , der erste Schmiegrum J_{12} von (71) durch die Vektoren dx_i und $\vartheta dx_i = \vartheta_2 x_i$.

Aus (70') folgt, daß:

$$(72) \quad J_{12} = \{A dx_i + B \vartheta_2 x_i\} = \{\overline{A} dx_i + \overline{B} d^2 x_i\}^7) \text{ ist.}$$

Der zweite Schmiegrum J_{123} von (71) ist aufgespannt durch die Vektoren dx_i , $\vartheta_2 x_i$, $\vartheta(\vartheta_2 x_i) = \vartheta_3 x_i$, also zufolge (72) und (70') ist

$$(73) \quad J_{123} = \{A dx_i + B \vartheta_2 x_i + C \vartheta_3 x_i\} = \{\overline{A} dx_i + \overline{B} d^2 x_i + \overline{C} d^3 x_i\}.$$

Allgemein ist der k^{te} Schmiegrum von (71) $J_{12 \dots k, k+1}$:

$$(74) \quad J_{12 \dots k, k+1} = \{\overline{A} dx_i + \overline{B} d^2 x_i + \dots + \overline{L} d^{(k+1)} x_i\}^7).$$

Aus (72), (73), (74) und aus der Definition des k^{ten} Schmiegrumes einer F_l im S_n :

$$(75) \quad x_i = x_i(y_1 y_2 \dots y_l)$$

folgt, daß der k^{te} Schmiegrum $J_{12 \dots k, k+1}$ der F_l :

⁶⁾ Ähnliche Probleme behandelt L. P. Einsenhardt: Non-Riemannien Geometry im III. Kapitel.

⁷⁾ Die Ausdrücke $d^2 x_i$, $d^3 x_i$, ... haben keinen geometrischen Sinn in unserem S_n , wohl aber die Mannigfaltigkeiten $\{\overline{A} dx_i + \overline{B} d^2 x_i\}$, $\{\overline{A} dx_i + \overline{B} d^2 x_i + \dots + \overline{L} d^{k+1} x_i\}$. $J_{12 \dots k+1}$ fällt in jedem Punkte der Kurve (71) mit der $k+1$ -dimensionalen Hyperebene E_{k+1} zusammen, welche in diesen Punkten von (71) die $k+1$ -dimensionalen Schmiegrumhyperebenen darstellt.

$$(76) \quad J_{12} \dots k_{k+1} \equiv \left\{ A_r \frac{\partial x_i}{\partial y_r} + A_{rs} \frac{\partial^2 x_i}{\partial y_r \partial y_s} + \dots + \right. \\ \left. + A_{r_1 r_2 \dots r_{k+1}} \frac{\partial^{(k+1)} x_i}{\partial y_{r_1} \dots \partial y_{r_{k+1}}} \right\}^s$$

ist.

Es sei nun ein n -dimensionaler topologischer Koordinatenraum T_n mit einem festen Koordinatensystem $x_1 x_2 \dots x_n$ gegeben. Jedem Punkte des Raumes T_n ordnen wir zu der Gesamtheit der Maßtensoren $\mathfrak{D} \{g_{ik}, \overline{g_{ik}}, \overline{\overline{g_{ik}}} \dots\}$ von der Art, daß die Geraden für die Maßtensoren $\mathfrak{D} \{g_{ik}, \overline{g_{ik}} \dots\}$ in diesem T_n für das Koordinatensystem $x_1, x_2, \dots x_n$ die Form

$$(61) \quad x_i = a_i t + b_i$$

haben.

Es sei nun in unserem T_n eine F_l :

$$(77) \quad x_i = x_i(y_1 y_2 \dots y_l)$$

gegeben, dann wollen wir die Invarianten (geometrische Eigenschaften) der F_l im T_n in bezug auf die Gesamtheit der Maßtensoren $\mathfrak{D} \{g_{ik}, \overline{g_{ik}}, \overline{\overline{g_{ik}}} \dots\}$ untersuchen. Diese Invarianten wollen wir kurz mit J_q bezeichnen.

Aus (76) folgt, daß der Raum

$$(78) \quad J_{12} \dots k_{k+1} \text{ eine } J_q \text{ ist.}$$

Bezeichnen wir die k^{ten} Normalräume von (77) in bezug auf die Tensoren $\mathfrak{D} \{g_{ik}, \overline{g_{ik}}, \overline{\overline{g_{ik}}} \dots\}$ mit $J_{h+1}, \overline{J_{k+1}}, \overline{\overline{J_{k+1}}}$, so folgt aus der Definition der Räume $J_{k+1}, \overline{J_{k+1}}, \overline{\overline{J_{k+1}}}, \dots$ und (78), daß die Dimensionalität der $J_{k+1}, \overline{J_{k+1}}, \overline{\overline{J_{k+1}}}, \dots$ eine J_q ist.

Bezeichnen wir jetzt die Riccische Ableitung des Vektors λ^i in bezug auf die Tensoren $g_{ik}, \overline{g_{ik}}, \overline{\overline{g_{ik}}}, \dots$ mit $\left(\frac{\partial \lambda^i}{\partial y_r}\right)_R, \left(\frac{\partial \lambda^i}{\partial y_r}\right)_{\overline{R}}, \dots$, so sieht man ohne weiteres, daß $\left(\frac{\partial \lambda^i}{\partial y_r}\right)_R, \left(\frac{\partial \lambda^i}{\partial y_r}\right)_{\overline{R}}, \dots$ Vektoren im T_n sind. Schreiben wir jetzt kurz Riccischen Ableitungen $\frac{\partial}{\partial y_{n_h}}, \left(\frac{\partial}{\partial y_{n_{h-1}}} \dots \left(\frac{\partial x_i}{\partial y_{n_1}}\right)_R\right)_R, \frac{\partial}{\partial y_{n_h}}, \left(\frac{\partial}{\partial y_{n_{h-1}}} \dots \left(\frac{\partial x_i}{\partial y_{n_1}}\right)_{\overline{R}}\right)_{\overline{R}}, \dots$

⁹⁾ Für die Ausdrücke $\frac{\partial^k x_i}{\partial y_{r_1} \partial y_{r_2} \dots \partial y_{r_k}}$ gilt dasselbe wie für $d^k x_i$.

$$\text{mit } \left(\frac{\partial^h x_i}{\partial y_{n_h} \dots \partial y_{n_1}} \right)_R, \left(\frac{\partial^h x_i}{\partial y_{n_h} \dots \partial y_{n_1}} \right)_{\bar{R}}, \dots$$

so sieht man, daß zufolge (70)

$$(79') \quad \left(\frac{\partial^h x_i}{\partial y_{n_1} \partial y_{n_2} \dots \partial y_{n_h}} \right)_R = \frac{\partial^h x_i}{\partial y_{n_1} \dots \partial y_{n_h}} + (.) \text{ und}$$

$$(79'') \quad \left(\frac{\partial^h x_i}{\partial y_{n_1} \partial y_{n_2} \dots \partial y_{n_h}} \right)_{\bar{R}} = \frac{\partial^h x_i}{\partial y_{n_1} \dots \partial y_{n_h}} + (.)^9.$$

Wir zeigen jetzt, daß, wenn für den Tensor g_{ik} der Gesamtheit $\mathfrak{Q}$ für die F_l die Beziehung:

$$(80) \quad T^{n_k \dots n_1} \left(\frac{\partial^k x_i}{\partial y_{n_1} \dots \partial y_{n_k}} \right)_R + (.)_R = 0$$

gilt, daß dann auch für jeden Tensor $\bar{g}_{ik}$ der Gesamtheit $\mathfrak{Q}$ die Beziehung

$$(81) \quad T^{n_k \dots n_1} \left(\frac{\partial^k x_i}{\partial y_{n_1} \dots \partial y_{n_k}} \right)_{\bar{R}} + (.)_{\bar{R}} = 0 \text{ gilt}^{10}.$$

Beweis:

Aus (80) und (79') folgt, daß:

$$(82) \quad 0 = T^{n_k \dots n_1} \left(\frac{\partial^k x_i}{\partial y_{n_1} \dots \partial y_{n_k}} \right)_R + (.) = T^{n_k \dots n_1} \frac{\partial^k x_i}{\partial y_{n_1} \dots \partial y_{n_k}} + (.)$$

Aus (79'') folgt, daß:

$$(83) \quad T^{n_k \dots n_1} \left(\frac{\partial^k x_i}{\partial y_{n_1} \dots \partial y_{n_k}} \right)_{\bar{R}} = T^{n_k \dots n_1} \frac{\partial^k x_i}{\partial y_{n_1} \dots \partial y_{n_k}} + (.)$$

⁹⁾ $(.)_R$ bedeutet, daß in $(.)_R$ nur Glieder $\left(\frac{\partial^r x_i}{\partial y_{n_1} \dots \partial y_{n_r}} \right)_R$ für $r \leq h-1$ vorkommen, analoges gilt für $(.)_{\bar{R}}$, während der Ausdruck $(.)$ bedeutet, daß in $(.)$ nur Glieder $\frac{\partial^r x_i}{\partial y_{n_1} \dots \partial y_{n_r}}$ für $r \leq h-1$ vorkommen.

¹⁰⁾ d. h.: Die Beziehung (81) ist eine J_q der Fläche F_l . Wir können dies auch folgendermaßen ausdrücken: Der Vektor $T^{n_k \dots n_1} \left(\frac{\partial^k x_i}{\partial y_{n_1} \dots \partial y_{n_k}} \right)_R$ liegt im $J_{1,2 \dots k-1}$ Raume für die Gesamtheit der Maßtensoren $\mathfrak{Q} \{g_{ik}, \bar{g}_{ik} \dots\}$, wenn der Vektor $T^{n_k \dots n_1} \left(\frac{\partial^k x_i}{\partial y_{n_1} \dots \partial y_{n_k}} \right)_R$ für einen Maßtensor g_{ik} der Gesamtheit $\mathfrak{Q}$ im $J_{1,2 \dots k-1}$ Raume liegt.

Subtrahiert man (82) von (83), so bekommt man:

$$(84) \quad T^{n_k \dots n_1} \left(\frac{\partial^k x_i}{\partial y_{n_k} \dots \partial y_{n_1}} \right)_{\bar{R}} = (.).$$

Da die linke Seite von (84) ein Vektor im T_n ist, so muß auch die rechte Seite von (84) ein Vektor im T_n sein, also auch im metrischen Raume mit dem Tensor $\bar{g}_{ik}$; da aber $(.)$ nur Glieder $\frac{\partial^h x_i}{\partial y_{n_h} \dots \partial y_{n_1}}$ für $h \leq k-1$ enthält, so läßt sich der Vektor $(.)$ in (84) durch die Vektoren $\left(\frac{\partial^h x_i}{\partial y_{n_h} \dots \partial y_{n_1}} \right)_{\bar{R}}$, für $h \leq k-1$ linear darstellen¹¹⁾, d. h. $(.) \equiv (.)_{\bar{R}}$. Es ist also:

$$(85) \quad T^{n_k \dots n_1} \left(\frac{\partial^k x_i}{\partial y_{n_k} \dots \partial y_{n_1}} \right)_{\bar{R}} + (.)_{\bar{R}} = 0$$

w. z. b. w.

Aus (79'), (79'') folgt, daß:

$$(86) \quad \left(\frac{\partial^k x_i}{\partial y_{n_k} \dots \partial y_{n_1}} \right)_R - \left(\frac{\partial^k x_i}{\partial y_{n_k} \dots \partial y_{n_1}} \right)_{\bar{R}} = (.) \text{ ist.}$$

Die letzte Gleichung (86) wollen wir anwenden auf den Fall $k=2$ und $J_2, \bar{J}_2 \dots$ eindimensional. Es bestehen dann die Gleichungen:

$$(87') \quad \left(\frac{\partial^2 x_i}{\partial y_p \partial y_q} \right)_R = B_{pq} \xi^i + (.) \frac{\partial x_i}{\partial y_i}$$

$$(87'') \quad \left(\frac{\partial^2 x_i}{\partial y_p \partial y_q} \right)_{\bar{R}} = \bar{B}_{pq} \bar{\xi}^i + (.) \frac{\partial x_i}{\partial y_i}$$

wobei $\xi^i, \bar{\xi}^i$ die die Räume $J_2, \bar{J}_2$ aufspannenden Einheitsvektoren sind. Aus (87'), (87'') folgt zufolge (86):

$$(88') \quad B_{pq} \xi^i - \bar{B}_{pq} \bar{\xi}^i = (.) \frac{\partial x_i}{\partial y_i}$$

entsprechend

$$(88'') \quad B_{rs} \xi^i - \bar{B}_{rs} \bar{\xi}^i = (.) \frac{\partial x_i}{\partial y_i}.$$

¹¹⁾ d. h. durch die Vektoren des $J_{1,2} \dots k-1$ Raumes.

Eliminiert man aus (88'), (88'') ξ^i , so bekommt man:

$$(89) \quad (B_{pq} \bar{B}_{rs} - B_{rs} \bar{B}_{pq}) \xi^i = (.) \frac{\partial x_i}{\partial y_t}.$$

Aus (89) folgt, da ξ^i ein Vektor des J_2 und $\frac{\partial x_i}{\partial y_t}$ Vektoren des J_1 sind, daß:

$$(90) \quad B_{pq} \bar{B}_{rs} - B_{rs} \bar{B}_{pq} = 0, \text{ d. h. } \bar{B}_{pq} = \alpha B_{pq} \text{ ist.}$$

Verschwundet demnach die Krümmung

$$(91') \quad B_{pq} B_{rs} - B_{ps} B_{rq} = 0,$$

so verschwindet aus zufolge (90)

$$(91'') \quad \bar{B}_{pq} \bar{B}_{rs} - \bar{B}_{ps} \bar{B}_{rq} = 0.$$

Da Verschwinden der Krümmung $B_{pq} B_{rs} - B_{ps} B_{rq} = 0$ ist im Falle, daß $J_2, \bar{J}_2 \dots$ eindimensional ist, eine Jq .

Wir haben somit den Satz:

Die F_l , deren J_2 Raum eindimensional ist und deren Krümmung (91') verschwindet, wird in einem ausgezeichneten Koordinatensystem (61) des S_n durch ein System von Gleichungen

$$(91''') \quad x_i = x_i(y_1, y_2, \dots, y_l)$$

gegeben, das im euklidischen Raume R_n für ein kartesisches Koordinatensystem die Hypertorse resp. die Hypertorse im erweiterten Sinne darstellt¹²⁾.

1. Bemerkung: Die Einbettungszahl k einer F_l im S_n ist¹³⁾:

$$(92) \quad k = \sum_{\sigma=1}^r l_{\sigma}, \text{ wo } l_{\sigma} \text{ die Dimension des } \sigma^{\text{ten}} \text{ Normalraumes } \bar{J}_{\sigma} \text{ der } F_l$$

ist. In der Tat sei:

¹²⁾ Denn die Gleichungen (37), (46), (53') kann man zufolge (72), (73), (74) in einer Jq -Form schreiben, z. B. (37) $(A): x_i = \bar{A}_i + (\bar{A}_i^{(r)})_R u_{r+1}$, wobei $(A^{(r)})_R$ die r te Riccische Ableitung von $\bar{A}_i$ in bezug auf den Tensor g_{ik} der Gesamtheit $\mathfrak{Q}$ ist. Aus der letzten Form (A) sieht man ohne weiteres die Jq -Eigenschaft von (A) . Die Form (A) ist aber zufolge (72), (73), (74) äquivalent der Form $(B): x_i = \bar{A}_i + \bar{A}_i^r u_{r+1}$. Die Flächen (A) , (B) sind demnach Jq und Lösungen unseres Problems für die Gesamtheit der Tensoren $\mathfrak{Q} \{g_{ik}, g_{ik}, \dots\}$, sie stellen also Hypertorsen resp. Hypertorsen im erweiterten Sinne im S_n dar.

¹³⁾ Siehe: Über das vollständige Formensystem der F_e im R_n von W. Mayer in Wien, Monatshefte für Math. und Physik, XXXV, Abschnitt II, § 2.

$$(92') \quad x_i = x_i (y_1 y_2 \dots y_l),$$

die F_l im S_n in einem ausgezeichneten Koordinatensystem (61) und seien $\bar{J}_\sigma$ die Normalräume der F_l . Da die Zahlen l_σ : J_q Eigenschaften sind, so haben die J_σ -Räume der F_l (92') im R_n mit Maßtensor $\delta_{ik} = \begin{matrix} 0, & i \neq k \\ 1, & i = k \end{matrix}$ auch die Dimension l_σ ; die Einbettungszahl der F_l (92') im R_n ist demnach gleich k , sie liegt also in einer E_k (k -dimensionaler) Hyperebene. Da aber jede E_k des R_n zugleich eine k -dimensionale Hyperebene $\bar{E}_k$ im ausgezeichneten Koordinatensystem (61) ist, so liegt die F_l (92') in E_k , d. h. ihre Einbettungszahl ist höchstens k . Auf dieselbe Art zeigt man, daß ihre Einbettungszahl mindestens k ist, demnach ist sie wirklich gleich k , w. z. b. w.

2. Bemerkung: Alle J_q -Invarianten der F_l sind zugleich projektive Invarianten der F_l .

III. Fortsetzung.

I. Satz: Ist für eine F_l im S_n der J_k -Raum eindimensional und existiert der J_{k+1} -Raum, dann verschwindet der $k-1^{\text{te}}$ Krümmungstensor.

$$(93) \quad A_{p_1 p_2 \dots p_k} \cdot A_{q_1 q_2 \dots q_k} - A_{p_1 p_2 \dots q_k} A_{q_1 q_2 \dots p_k} = 0.$$

Beweis: Bezeichnen wir $A_{p_1 p_2 \dots p_k}$ mit ${}^{(\alpha)}A_{p_1 p_2 \dots p_k}$, dann ist

$$(94) \quad {}^{(\alpha)}A_{p_1 p_2 \dots p_k} {}^{(\alpha \beta)}C_{p_{k+1}} = {}^{(\beta)}A_{p_1 p_2 \dots p_k p_{k+1}} {}^{14)},$$

wobei β ein Index des J_{k+1} Raumes ist.

Aus der Symmetrie von ${}^{(\beta)}A_{p_1 p_2 \dots p_k p_{k+1}}$ in den Indizes¹⁵⁾ $p_1 p_2 \dots p_k p_{k+1}$ und aus (94) folgt:

$$(95) \quad {}^{(\beta)}A_{p_1 p_2 \dots p_{k+1}} = {}^{(\beta)}B C_{p_1} C_{p_2} \dots C_{p_{k+1}}, \quad 9d$$

wobei wir für ${}^{(\alpha \beta)}C_{p_\alpha}$ kurz C_{p_α} gesetzt haben. 30)

Aus (94), (95) folgt dann:

$$(96) \quad {}^{(\alpha)}A_{p_1 p_2 \dots p_k} = {}^{(\beta)}B C_{p_1} C_{p_2} \dots C_{p_k}.$$

Aus (96) folgt dann das identische Verschwinden des Tensors (93), w. z. b. w.

Bemerkung: Aus (95) folgt, daß auch der J_{k+1} eindimensional ist und daß der k^{te} Krümmungstensor:

¹⁴⁾ Siehe Formenproblem, I. Kap., § 3, II. Kap., § 4.

¹⁵⁾ Siehe Formenproblem, I. Kap., § 12, II. Kap., § 4.

$$(97) \quad A_{p_1 p_2 \dots p_k p_{k+1}} A_{q_1 q_2 \dots q_k q_{k+1}} - A_{p_1 p_2 \dots p_k q_{k+1}} A_{q_1 q_2 \dots q_k p_{k+1}} = 0$$

identisch verschwindet.

II. Satz: Ist für eine Regelhyperfläche F_l im S_n der erste Krümmungstensor

$$(98) \quad {}_{(\alpha)}B_{pq} {}_{(\alpha)}B_{rs} - {}_{(\alpha)}B_{ps} {}_{(\alpha)}B_{rq} = 0$$

dann ist der J_2 Raum der F_l eindimensional, die F_l ist daher eine Hypertorse.

Beweis: Für das ausgezeichnete Koordinatensystem (61) können wir nach Voraussetzung die F_l darstellen in der Form:

$$(99) \quad x_i = A_i(y_1) + C_{ik}(y_1) \cdot y_k; \quad k = 2, 3 \dots l; \quad i = 1, 2 \dots n$$

und es gelten die Beziehungen:

$$(100) \quad \left(\frac{\partial^2 x_i}{\partial y_p \partial y_q} \right)_R = {}_{(\alpha)}B_{pq} {}_{(\alpha)}\lambda^i{}_{16},$$

wo α die Indizes des J_2 Raumes durchläuft.

Für die F_l (99) ist

$$(101) \quad \left(\frac{\partial^2 x_i}{\partial y_p \partial y_q} \right)_R = 0 \quad \text{wenn } p, q = 2, 3, \dots l \text{ ist.}$$

Wäre nun der J_2 Raum nicht eindimensional, so müßte für mindestens ein $k \neq 1$

$$(102) \quad \left(\frac{\partial^2 x_i}{\partial y_1 \partial y_k} \right)_R = {}_{(\alpha)}B_{1k} {}_{(\alpha)}\lambda^i \neq 0 \text{ sein,}$$

da nach (101) alle ${}_{(\alpha)}B_{rs} = 0$ sind, für $r, s = 2, 3 \dots l$. Dann sind aber nicht alle ${}_{(\alpha)}B_{1k}$ Null, somit ist

$$(103) \quad \sum_{\alpha} ({}_{(\alpha)}B_{1k})^2 \neq 0.$$

Nun bilden wir:

$$(104) \quad {}_{(\alpha)}B_{11} {}_{(\alpha)}B_{kk} - {}_{(\alpha)}B_{1k} {}_{(\alpha)}B_{1k},$$

so ist (104) zufolge (98) gleich Null, andererseits ist ${}_{(\alpha)}B_{kk} = 0$ zufolge (101) und $\sum_{\alpha} {}_{(\alpha)}B_{1k}^2 \neq 0$ zufolge (103), dies ergibt demnach einen Widerspruch. Somit ist der Satz II bewiesen.

¹⁶⁾ Siehe Formenproblem, I. Kap., § 3, II. Kap., § 4.

Wir wollen jetzt an einem Beispiele zeigen, daß die Tatsache des Verschwindens der ersten Krümmung der F_l in einem R_n ¹⁷⁾ nicht beweist, daß eine F_l eine Hypertorse resp. eine Hypertorse im weiteren Sinne sei. Wesentlich tritt noch die Eindimensionalität des J_2 Raumes als Voraussetzung hinzu.

Wir führen das Beispiel für den R_n an, dem wir ein orthogonales kartesisches Koordinatensystem mit $g_{ik} = \delta_{ik}$ als Maßtensor zugrunde legen.

Wir bilden aus den Zahlen $1, 2, \dots, n$, die l Gruppen

$$(105) \quad \left\{ \begin{array}{l} k_1: 1, 2, \dots, k_1 \\ k_2: k_1 + 1, k_1 + 2, \dots, k_2 \\ \vdots \\ k_l: k_{l-1} + 1, \dots, k_l = n, \end{array} \right.$$

wobei wir voraussetzen, daß zumindest zwei Gruppen (105) z. B. k_p und k_q mehr als eine Zahl enthalten.

Wir definieren nun die F_l durch die Gleichungen:

$$(106) \quad \begin{array}{l} x_i = \varphi_{i1}(y_1), \quad i = 1, 2, \dots, k_1 \\ x_i = \varphi_{i2}(y_2), \quad i = k_1 + 1, \dots, k_2 \\ \vdots \\ x_i = \varphi_{il}(y_l), \quad i = k_{l-1} + 1, \dots, k_l \end{array}$$

und treffen die Vereinbarung, daß die Kurven

$$(107) \quad \left\{ \begin{array}{l} x_i = \varphi_{ir}(y_r) \\ x_i = 0 \end{array} \right.$$

(für alle Zahlen i von k_t , wenn $t \neq r$ ist), für $r = p$ und $r = q$ keine Geraden sind. Dies bedeutet, daß für die Kurven (107) der Vektor

$$(108) \quad \frac{d_{(1)} \xi_i}{ds} \neq 0 \text{ ist, für } r = p \text{ und } r = q.$$

Wir bilden nun für die F_l (106) den Maßtensor γ_{rt} :

$$(109) \quad \gamma_{rt} = \frac{\partial x_i}{\partial y_r} \cdot \frac{\partial x_i}{\partial y_t},$$

¹⁷⁾ Siehe Formenproblem, I. Kap., § 11, Formel (9).

aus (106) folgt dann, daß

$$(110) \quad \gamma_{rt} = 0, \text{ für } r \neq t$$

und daß:

$$(111) \quad \gamma_{rr} = \frac{\partial \varphi_{ir}(y_r)}{\partial y_r} \frac{\partial \varphi_{ir}(y_r)}{\partial y_r} = \Phi_r(y_r)$$

eine Funktion von y_r allein ist. Somit ist das Bogenelement ds^2 von F_l (106) von der Form:

$$(112) \quad ds^2 = \sum_{r=1}^l \Phi_r(y_r) dy_r^2.$$

Durch die Transformation

$$(113) \quad z_p = \int \sqrt{\Phi_r(y_r)} dy_r$$

erhält man für ds^2 die Form:

$$(114) \quad ds^2 = \sum_{r=1}^l dz_r^2.$$

Demnach verschwinden die Klammern $\left\{ \begin{smallmatrix} p & q \\ r \end{smallmatrix} \right\}$ und die erste Krümmung der F_l (106) ist Null. Es bleibt noch zu zeigen, daß der J_2 Raum mehr als eindimensional ist.

Für die zweiten Ableitungen gilt die Darstellung:

$$(115) \quad \frac{\partial^2 x_i}{\partial y_r \partial y_t} = \left\{ \begin{smallmatrix} r & t \\ j \end{smallmatrix} \right\} \frac{\partial x_i}{\partial y_j} + \frac{\partial^2 x_i}{\partial y_r \partial y_t} \quad \text{und}$$

$$(115') \quad \frac{\partial^2 x_i}{\partial z_r \partial z_t} = \frac{\partial^2 x_i}{\partial z_r \partial z_t},$$

da ja $\left\{ \begin{smallmatrix} r & t \\ j \end{smallmatrix} \right\}$ für Parametersystem z_r (113), (114) verschwindet. Aus (110), (111), (113) folgt, daß

$$(116) \quad \frac{\partial^2 x_i}{\partial x_r \partial z_t} = 0, \text{ für } r \neq t \text{ ist, also } \frac{\partial^2 x_i}{\partial z_r \partial z_t} = 0 \text{ ist.}$$

Von den Vektoren des J_2 Raumes können daher nur die Vektoren

$$(117) \quad \frac{\partial^2 x_i}{\partial z_r \partial z_r}$$

von Null verschieden sein.

Wir zeigen, daß die zwei Vektoren (117) $r=p$, $r=q$ nicht verschwinden.

Für die Bogenlänge der Kurven (107) gilt:

$$(118) \quad ds^2 = \sum_{i=k_{p-1}+1}^{k_r} dx_i^2 = \sum_i \frac{\partial x_i}{\partial y_r} \frac{\partial x_i}{\partial y_r} dy_r^2 = \Phi_r dy_r^2 = dz_r^2$$

(über r wird nicht summiert), somit ist z_r die Bogenlänge der Kurve (110) für $r=p$.

Dann folgt für diese Kurve, für $i = k_{p-1} + 1, \dots, k_p$

$$(119) \quad {}_{(1)}\xi_i = \frac{dx_i}{dz_p},$$

wogegen ${}_{(1)}\xi_i = 0$ ist, wenn i keine Zahl der Gruppe k_p bedeutet. Damit gilt

$$(120) \quad \frac{d {}_{(1)}\xi_i}{ds} = \frac{d {}_{(1)}\xi_i}{dz_p} = \frac{d^2 x_i}{(dz_p)^2} = \frac{d^2 x_i}{\underline{(dz_p)^2}}$$

zufolge (115'), also zufolge (108) ist $\frac{\partial^2 x_i}{\partial z_p \partial z_p} \neq 0$.

Dasselbe gilt für $\frac{\partial^2 x_i}{\partial z_q \partial z_q} \neq 0$.

Wäre nun der J_2 Raum eindimensional, so müßten die folgenden Gleichungen zusammen bestehen:

$$(121) \quad B_{pq} = 0, \quad B_{pp} \neq 0, \quad B_{qq} \neq 0$$

$$B_{pp} \cdot B_{qq} - (B_{pq})^2 = 0$$

[zufolge (115'), (120)]. Dies ist aber unmöglich, also ist der J_2 Raum der F_l (106) mehr als eindimensional, w. z. b. w.

IV. Über eine spezielle Klasse Riemannscher Räume.

Es sei gegeben ein Riemannscher Raum V_n mit einem symmetrischen nicht ausgearbeiteten Maßtensor g_{ik} . Die Parallelverschiebung der Vektoren $\lambda_i \dots \mu_i \dots$ ist dann gegeben durch das System

$$(122') \quad d\lambda_i = - \left\{ \begin{matrix} pq \\ i \end{matrix} \right\} \lambda^q dx_q,$$

$$(122'') \quad d\mu_i = \left\{ \begin{matrix} i p \\ q \end{matrix} \right\} \mu_p dx_q^{18}.$$

Die Parallelverschiebung läßt das innere Produkt $\lambda^i \mu_i$ invariant und zerstört nicht den Zusammenhang gekoppelter Vektoren¹⁹⁾. Der Maßtensor g_{ik} erfüllt die Identität

$$(123') \quad \frac{\partial g_{ik}}{\partial x_j} - \left\{ \begin{matrix} i j \\ l \end{matrix} \right\}_g g_{lk} - \left\{ \begin{matrix} k j \\ l \end{matrix} \right\}_g g_{li} = 0.$$

Gilt für einen symmetrischen, nicht ausgearbeiteten Tensor a_{ik}

$$(123'') \quad \frac{\partial a_{ik}}{\partial x_j} - \left\{ \begin{matrix} i j \\ l \end{matrix} \right\}_g a_{lk} - \left\{ \begin{matrix} k j \\ l \end{matrix} \right\}_g a_{li} = 0,$$

so ist dann

$$(124') \quad \left\{ \begin{matrix} i k \\ l \end{matrix} \right\}_g = \left\{ \begin{matrix} i k \\ l \end{matrix} \right\}_a$$

und der Ausdruck

$$(124'') \quad a_{ik} \lambda^i \mu^k$$

bleibt bei Parallelverschiebung der Vektoren λ^i , μ^i konstant, wie man ohne weiters aus (122'), (122''), (123'), (123'') einsehen kann²⁰⁾.

Nach diesen Vorbereitungen fragen wir nach den Riemannschen Räumen V_n , in denen das System (123') eine von g_{ik} (resp. $C \cdot g_{ik}$, $C = \text{konstant}$) verschiedene Lösung zuläßt, in denen also der Maßtensor g_{ik} durch die Parallelverschiebung nicht eindeutig bestimmt ist.

Gibt es einen solchen Tensor a_{ik} , so haben wir gesehen, daß der Ausdruck (124'') für parallel verschobene Vektoren λ^i , μ_i konstant bleibt. Wir betrachten dann in jedem Punkte P des V_n die Gesamtheit der Einheitsvektoren λ^i (Richtungen):

$$(125) \quad g_{ik} \lambda^i \lambda^k = 1$$

und suchen jene Richtungen, für die

$$(126) \quad a_{jk} \lambda^j \lambda^k = \sigma$$

ein Extremum wird.

¹⁸⁾ $\left\{ \begin{matrix} i k \\ l \end{matrix} \right\}_g$ sind die Christoffelklammern des Tensors g_{ik} entsprechend $\left\{ \begin{matrix} i k \\ l \end{matrix} \right\}_a$ des Tensor a_{ik} .

¹⁹⁾ Es ist $\lambda_i = g_{ik} \lambda^k$, $\lambda^i = g^{ik} \lambda_k$.

²⁰⁾ Der Tensor a_{ik} kann in diesem Falle als ein Maßtensor und als Kopplungstensor im V_n verwendet werden.

Die Extremalrichtungen genügen dem System linearer Gleichungen:

$$(127) \quad a_{ik} \lambda^i - \rho g_{ik} \lambda^i = 0,$$

wobei wegen (126):

$$a_{ik} \lambda^i \lambda^k = \rho = \sigma \text{ ist.}$$

Die n Wurzeln ρ der Gleichung:

$$(128) \quad |a_{ik} - \rho g_{ik}| = 0,$$

die bekanntlich reell sind, stellen die Extremalwerte von (126) dar²¹⁾.

Sei nun: ρ_1 eine l_1 — fache Wurzel von (128)

$$\begin{array}{ccccccc} \rho_2 & n & l_2 & — & & & \\ \rho_t & n & l_t & — & & & \end{array}$$

wobei $\rho_1, \rho_2 \dots \rho_t$ sämtliche Wurzeln von (128) sind, also $l_1 + l_2 + \dots + l_t = n$ ist. Bezeichnen wir mit W_h den l_h dimensionalen Vektorraum der Extremalrichtungen λ^i , für die $a_{ik} \lambda^i \lambda^k = \rho_h$ ist, so enthält der n dimensionale Vektorraum in jedem Punkte P des V_n die Vektorräume der Extremalrichtungen: $W_1, W_2, \dots, W_t$, die zu je zweien aufeinander normal stehen.

Wir fixiren nun im V_n zwei Punkte P_1 und P_2 , die wir uns durch beliebige Kurven verbunden denken. Wenn wir nun längs dieser Kurven die Richtungen λ^i von P_1 nach P_2 parallel verschoben denken, so bleiben für jede dieser Richtungen die Gleichungen (125), (126) bestehen. Die Extremalrichtungen für den Wert ρ_h in P_1 gehen in Extremalrichtungen für diesen Wert in P_2 über. Also führt die Parallelverschiebung die Räume W_h von P_1 über in die Räume W_h in P_2 , und zwar unabhängig von den P_1 mit P_2 verbindenden Kurven.

Im V_n gibt es demnach Vektorräume W_h , die bei Parallelverschiebung erhalten bleiben.

²¹⁾ Sind ρ_1 und ρ_2 zwei ungleiche Wurzeln von (128) und λ^i, λ^i , die entsprechenden Extremalrichtungen, so folgt aus (127): (1) (2)

$$\begin{array}{ccc} a_{ik} \lambda^i \lambda^k = \rho_1 g_{ik} \lambda^i \lambda^i = \rho_2 g_{ik} \lambda^i \lambda^k. \\ (1) (2) \qquad (1) (2) \qquad (1) (2) \end{array}$$

Und da $\rho_1 \neq \rho_2$ ist, so folgt daraus $a_{ik} \lambda^i \lambda^k = g_{ik} \lambda^i \lambda^k = 0$, d. h. Extremalrichtungen, die ungleichen Wurzeln von (128) entsprechen, sind aufeinander senkrecht. Ist ρ_v eine l_v fache Wurzel von (128), so ist der Rang der Determinante (128) $|a_{ik} - \rho_v g_{ik}| = 0$, $n - l_v$. Es gibt demnach l_v linear unabhängige Extremalrichtungen für ρ_v (jede andere Extremalrichtung läßt sich durch diese l_v Richtungen linear ausdrücken). Die Extremalrichtungen für die Wurzel ρ_v spannen also einen l_v dimensionalen Vektorraum aus.

Wir behaupten nun auch umgekehrt:

Gibt es im V_n Vektorräume W_h , die invariant sind in bezug auf die Parallelverschiebung, so ist der Maßtensor g_{ik} von V_n nicht eindeutig bestimmt²²⁾.

Seien $W_1, W_2, \dots, W_t$ die Vektorräume, die zu je zweien normal stehen, so wählen wir ein kontravariantes n -Bein²³⁾.

$$(129) \quad {}_{(\alpha)}\lambda^i; \alpha, i = 1, 2, \dots, n,$$

so daß die Vektoren

$$(130) \quad {}_{(\alpha_h)}\lambda^i; \alpha_h = l_1 + l_2 + \dots + l_{h-1} + 1, l_1 + l_2 + \dots + l_{h-1} + 2, \dots \\ \dots l_1 + l_2 + \dots + l_{h-1} + l_h,$$

den W_h Raum aufspannen²⁴⁾. Es ist dann:

$$(131) \quad g_{ik} = {}_{(\alpha)}\lambda_i {}_{(\alpha)}\lambda_k; g^{ik} = {}_{(\alpha)}\lambda^i {}_{(\alpha)}\lambda^k.$$

Gibt es nun einen g_{ik} verschiedenen Maßtensor, so folgt aus (127)

$$(132) \quad a_{ik} {}_{(\alpha_h)}\lambda^i = \rho_h g_{ik} {}_{(\alpha_h)}\lambda^i, \quad h = 1, 2, \dots, t \\ i, k = 1, 2, \dots, n,$$

wobei, wie wir zeigten, ρ_h im V_n konstante Größen sind.

Aus (131), (132) folgt:

$$(133) \quad a_{ik} = \sum_{h=1}^t \rho_h g_{ik} {}_{(\alpha_h)}\lambda^i {}_{(\alpha_h)}\lambda_k = \sum_{h=1}^t \rho_h {}_{(\alpha_h)}\lambda^i {}_{(\alpha_h)}\lambda_k \quad (\rho_h \text{ konstant})$$

und (133) ist die allgemeinste Form des a_{ik} dieser Eigenschaft. Umgekehrt entspricht ein solcher a_{ik} (133) mit beliebigen Konstanten ρ_h unserer Forderung.

In Riemannschen Räumen V_n , in denen es Vektorräume W_h gibt, die in bezug auf die Parallelverschiebung invariant sind, ist der Tensor g_{ik} nicht eindeutig fixiert.

Solche Räume sind in der Literatur bereits betrachtet worden und es wurde für sie gezeigt, daß man die invarianten Vektorräume W_h zu Hyperflächen verbinden kann, die Hyperebenen im V_n sind.

²²⁾ D. h. es gibt einen a_{ik} , wobei $a_{ik} = C g_{ik}$ (C konstant), für welchen (123'') gilt. Die Vektorräume W_h stehen dann zu je zweien aufeinander senkrecht.

²³⁾ D. h. es ist $g_{ik} {}_{(\alpha)}\lambda^i {}_{(\beta)}\lambda^k = \delta_{\alpha\beta}$.

²⁴⁾ Wir werden die Indizes, welche die Reihe $l_1 + l_2 + \dots + l_{h-1} + 1, l_1 + l_2 + \dots + l_{h-1} + 2, \dots, l_1 + l_2 + \dots + l_{h-1} + l_h$ durchlaufen, mit $\alpha_h, \dots, \rho_h, \dots, q_h$ bezeichnen.

Bezeichnen wir mit E_{l_h} die l_h dimensionalen Hyperebenen, die die W_h enthalten und führen wir im V_n solche Koordinaten x_i ein, so daß die Gleichungen der E_{l_h} gegeben sind durch folgende Relationen:

$$(134) \quad E_{l_h}: x_i = c_r; \text{ für } r \neq \alpha_h \quad (\alpha_h = l_1 + l_2 + \dots + l_{h-1} + 1, l_1 + l_2 + \dots + l_{h-1} + 2, \dots, l_1 + l_2 + \dots + l_{h-1} + l_h)$$

(c_r konstant). Da die Hyperflächen E_{l_h} zueinander normal stehen, so ist:

$$(134') \quad g_{ik} = 0,$$

wenn i und k Indizes verschiedener W_h Räume sind. Da E_{l_h} Hyperebenen sind, so liegt jede Gerade $x_i = x_i(s)$ (s Bogenlänge) in ihr, wenn sie einen Punkt P und in P eine Richtung mit der E_{l_h} gemeinsam hat. Die Gleichungen für diese Gerade lauten:

$$(135) \quad d^2 x_i + \left\{ \begin{matrix} p & q \\ i & i \end{matrix} \right\} dx_p dx_q = 0$$

da $dx_i = d^2 x_i = 0$ ist, für $i \neq \alpha_h$ zufolge (130), (134) ist, so folgt aus (135)

$$(136') \quad 0 = \left\{ \begin{matrix} p_h & q_h \\ i & i \end{matrix} \right\} dx_{p_h} dx_{q_h}, \text{ für } i \neq \alpha_h \text{ also}$$

$$(136'') \quad \left\{ \begin{matrix} p_h & q_h \\ i & i \end{matrix} \right\} = 0, \text{ für } i \neq \alpha_h, \text{ wobei } p_h, q_h, \alpha_h$$

Indizes des W_h Raumes sind. Nun ist:

$$(137) \quad 0 = \left\{ \begin{matrix} p_h & q_h \\ i & i \end{matrix} \right\} = g^{ij} \left[\begin{matrix} p_h & q_h \\ j & j \end{matrix} \right] = \Sigma g^{ij} \left[\begin{matrix} p_h & q_h \\ j & j \end{matrix} \right]^{25}, \quad i \neq \alpha_h.$$

Da die Determinante $|g^{ij}|$ (wobei $i \neq \alpha_h, j \neq \alpha_h$ ist) von Null verschieden ist, somit folgt aus (137)

$$(138) \quad \left[\begin{matrix} p_h & q_h \\ j & j \end{matrix} \right] = 0, \text{ für } j \neq \alpha_h. \text{ Da aber}$$

$$(139) \quad 2 \left[\begin{matrix} p_h & q_h \\ j & j \end{matrix} \right] = \frac{\partial g_{p_h, j}}{\partial x_{q_h}} + \frac{\partial g_{j, q_h}}{\partial x_{p_h}} - \frac{\partial g_{p_h, q_h}}{\partial x_j} = - \frac{\partial g_{p_h, q_h}}{\partial x_j},$$

²⁵⁾ Aus (134') folgt, daß in (137) über die Indizes j des W_h Raumes nicht summiert wird.

für $j \neq \alpha_h$, zufolge (134), so ist

$$(140) \quad g_{p_h, q_h} = f_{p_h, q_h} (x_{\alpha_h})^{26},$$

es ist also

$$(141) \quad \begin{cases} g_{p_h, q_k} \equiv 0, & h \neq k \text{ und} \\ g_{p_h, q_h} = f_{p_h, q_h} (x_{\alpha_h}). \end{cases}$$

Umgekehrt folgt aus (141), daß für den Maßtensor (141) die Vektorräume W_{h1} bei der Parallelverschiebung invariant bleiben, dabei sind die W_h aufgespannt durch die Vektoren des n -Beines:

$$(142) \quad (x_h) \lambda^i = 0, \quad \text{für } i \neq \alpha_h, \quad h = 1, 2, \dots, t^{27}.$$

Aus (131) folgt, wenn wir (142) berücksichtigen, daß

$$(143) \quad g_{ik} = \sum_{h=1}^t (\alpha_h) \lambda_i (\alpha_h) \lambda_k$$

ist, also zufolge (133).

$$(144) \quad a_{ik} = \sum_{h=1}^t \rho_h (\alpha_h) \lambda_i (\alpha_h) \lambda_k;$$

aus (143) und (144) folgt:

$$(145) \quad a_{l_h k_h} = \rho_h g_{l_h k_h}, \quad \text{und} \quad a_{l_h k_h} = 0, \quad \text{für } h \neq r, \quad \text{wobei } \rho_h \text{ kon-}$$

stant ist. Man sieht ohne weiteres, daß

$$(145') \quad \left\{ \begin{smallmatrix} p & q \\ i \end{smallmatrix} \right\}_g = \left\{ \begin{smallmatrix} p & q \\ i \end{smallmatrix} \right\}_a \quad \text{ist und daß für den Tensor } a_{ik} \text{ (145) die}$$

²⁶⁾ d. h. $f_{p_h, q_h} (x_{\alpha_h})$ ist eine Funktion nur der Variablen x_{α_h} , wo α_h alle Indizes des W_h Raumes durchläuft.

²⁷⁾ Es ist nämlich $\left[\begin{smallmatrix} p & q \\ p & k \end{smallmatrix} \right] = 0$, wenn $h \neq k$ ist, also $\left\{ \begin{smallmatrix} p & q \\ p & k \end{smallmatrix} \right\}$, für $h \neq k$. Es sei $x_i = x_i(s)$, eine Kurve, längs der wir den Vektor $(\alpha_h) \lambda^i / P_1$ von P_1 aus parallel verschieben, dann haben wir das System

$$(\beta) \quad \frac{d(\alpha_h) \lambda^i}{ds} + \left\{ \begin{smallmatrix} p & q \\ i \end{smallmatrix} \right\} (\alpha_h) \lambda^p \frac{dx^q}{ds} = 0$$

zu integrieren. Wir können demnach (β) in der Form

$$(\gamma) \quad \frac{d(\alpha_h) \lambda^i}{ds} + \ddot{\Sigma} \left\{ \begin{smallmatrix} p & q \\ i \end{smallmatrix} \right\} (\alpha_h) \lambda^p \frac{dx^q}{ds} = 0$$

schreiben, wobei $\ddot{\Sigma}$ andeuten soll, daß nur über jene p summiert wird, die Indizes jenes W_h sind, dem i angehört. Dadurch zerfällt (γ) in zwei Systeme, je nachdem i ein Index von W_h ist oder nicht. Ist i kein Index des W_h Raumes, so können wir das entsprechende Teilsystem (γ) durch $(\alpha_h) \lambda^i = 0$ für $i \neq \alpha_h$ befriedigen, was dann dem Ausgangswerte von $(\alpha_h) \lambda^i$ im P_1 entspricht, und da nur diese Lösung die angegebenen Anfangswerte hat, so ist dadurch unsere Behauptung bewiesen.

Relation (123'') gilt. Somit haben wir den Satz bewiesen: Die notwendige und hinreichende Bedingung dafür, daß für einen Riemannschen Raum V_n bei gegebener Parallelverschiebung und für ein gegebenes Koordinationssystem der Maßtensor g_{ik} nicht eindeutig bestimmt ist, ist die Invarianz von Vektorräumen in bezug auf die Parallelverschiebung des V_n .

Bemerkung: Im Riemannschen Sinne sind die V_n Räume:

$$V_n \{x_i, g_{ik}\}, \quad V_n \{x_i, a_{ik}\}$$

identisch, denn durch die Transformation

$$(146) \quad \bar{x}_{i_h} = \sqrt{\rho_h} \ x_{i_h}$$

(über h nicht summiert) erhält man aus:

$$(147') \quad \bar{g}_{i_h j_h} \bar{d}x_{i_h} \bar{d}x_{j_h} = g_{i_h j_h} dx_{i_h} dx_{j_h}$$

$$(147'') \quad \bar{g}_{i_h j_h} \rho_h = g_{i_h j_h},$$

also zufolge (146)

$$(148) \quad g_{i_h j_h} = a_{i_h j_h}.$$

Im Koordinatensystem x_i ist demnach der transformierte Maßtensor $\bar{g}_{ik}$ gleich dem Tensor a_{ik} . Somit gilt der Satz:

Durch die Parallelverschiebung des Riemannschen Raumes ist dieser Raum völlig gegeben.

I. Nachtrag.

Über das System der Formenquadrate einer F_l im S_n .

Wir wollen hier in einer Richtung die Arbeit „Über das vollständige Formensystem einer F_l im R_n “ ergänzen, indem wir zeigen, daß die Formenquadrate:

$$(149) \quad \left\{ \begin{array}{l} \Sigma_{(\alpha)} I^2 = {}_{(\alpha)} A_p {}_{(\alpha)} A_q dy_p dy_q = B_{pq} dy_p dy_q \\ \Sigma_{(\alpha)} II^2 = {}_{(\alpha)} A_{pq} {}_{(\alpha)} A_{rs} dy_p \dots dy_s = B_{pqrs} dy_p \dots dy_s \end{array} \right\}^{(28)}$$

die F_l im S_n bis auf Kongruenz völlig festlegen.

²⁸⁾ Die geometrische Bedeutung der Formen (149) ist sehr einfach: Aus der Formel (49), II. Kap., § 4 (Formenproblem)

$${}_{(\alpha)} A_{p_1 p_2 \dots p_k} \frac{dy_{p_1}}{ds} \dots \frac{dy_{p_k}}{ds} = \frac{{}_{(\alpha)} \lambda_l {}_{(k)} \xi^l}{\rho_1 \rho_2 \dots \rho_{k-1}}$$

$${}_{(\alpha)} A_{p_1 p_2 \dots p_k} {}_{(\alpha)} A_{q_1 q_2 \dots q_k} \frac{dy_{p_1}}{ds} \dots \frac{dy_{q_k}}{ds} = \left(\frac{{}_{(k)} \xi_l {}_{(k)} \xi^l}{\rho_1 \rho_2 \dots \rho_{k-1}} \right)^2,$$

wobei ${}_{(k)} \xi^i$ die orthogonale Projektion von ${}_{(k)} \xi_i$ in den J_k Raum bedeutet.

Beweis: Sei eine F_l im S_n gegeben und ${}_{(\alpha)}A_{p_1 p_2 \dots p_k}$ ${}_{(\alpha)}\bar{A}_{p_1 p_2 \dots p_k}$ zwei äquivalente Formensysteme²⁹⁾, die den äquivalenten Beinen ${}_{(\alpha)}\lambda_i$ resp. ${}_{(\alpha)}\bar{\lambda}_i$ adjungiert seien³⁰⁾. Dann gilt

$$(150) \quad {}_{(\alpha)}\lambda_i = {}_{(\alpha\beta)}a_{(\beta)}\lambda_i,$$

(wobei α, β Indizes des J_k ($k=1, 2 \dots$) Raumes sind), wobei die Matrix $\| {}_{(\alpha\beta)}a \|$ eine orthogonale Matrix ist. Aus (150) folgt dann:

$$(151) \quad {}_{(\alpha)}\bar{A}_{p_1 p_2 \dots p_k} = {}_{(\alpha\beta)}a_{(\beta)}A_{p_1 p_2 \dots p_k}{}^{31)},$$

α, β Indizes des J_k Raumes als die notwendige und hinreichende Bedingung der Äquivalenz der beiden Formensysteme.

Aus (151) folgt weiter:

$$(152) \quad {}_{(\alpha)}\bar{A}_{p_1 p_2 \dots p_k} {}_{(\omega)}\bar{A}_{q_1 q_2 \dots q_k} = {}_{(\alpha)}A_{p_1 p_2 \dots p_k} {}_{(\omega)}A_{q_1 q_2 \dots q_k} = \\ = E_{p_1 p_2 \dots p_k} | q_1 q_2 \dots q_k$$

und daß die Tensoren

$$(152') \quad E_{p_1 p_2 \dots p_k} | q_1 q_2 \dots q_k \quad (k=1, 2 \dots)$$

die Formen ${}_{(\alpha)}A_{p_1 p_2 \dots p_k}$ bis auf Äquivalenz (151) vollständig bestimmen.

In der Tat, sei l_k die Anzahl der linear unabhängigen Tensoren ${}_{(\alpha)}\bar{A}_{p_1 p_2 \dots p_k}{}^{32)}$, dann gibt es $h_k = \binom{l+k-1}{k} - l_k$ Tensoren ${}_{(\alpha)}\bar{A}_{p_1 p_2 \dots p_k}{}^{33)}$, welche der Relation:

$$(153) \quad {}_{(\alpha)}\bar{A}_{q_1 q_2 \dots q_k} {}_{(\beta)}\bar{A}_{q_1 q_2 \dots q_k} = \delta_{\alpha\beta}$$

(α, β Indizes des J_k Raumes) genügen. Aus (152) folgt dann, indem man (152) mit ${}_{(\beta)}\bar{A}_{q_1 q_2 \dots q_k}$ multipliziert:

²⁹⁾ C. Burstin und W. Mayer: Das Formenproblem, I, Kap., § 3, II. Kap., § 4.

³⁰⁾ C. Burstin und W. Mayer: Das Formenproblem, I. Kap., § 3.

³¹⁾ C. Burstin und W. Mayer: Das Formenproblem, I. Kap., § 3.

³²⁾ Wobei ${}_{(\alpha)}\bar{A}_{p_1 p_2 \dots p_k}$ eine bestimmte Lösung des Systems (152) ist.

³³⁾ Es gibt ∞h_k Tensoren, welche der Relation (153) genügen, es sei ${}_{(\alpha)}\bar{A}_{p_1 p_2 \dots p_k}$ ein bestimmter Tensor dieser Art.

$$(154) \quad {}_{(\beta)}\bar{A}_{p_1 p_2 \dots p_k} = {}_{(\beta\alpha)}\bar{b}_{(\alpha)} A_{p_1 p_2 \dots p_k},$$

wobei

$$(154') \quad {}_{(\beta\alpha)}\bar{b} = {}_{(\alpha)}A_{q_1 q_2 \dots q_k} {}_{(\beta)}\bar{A}^{q_1 q_2 \dots q_k} \text{ ist.}$$

Wir beweisen nun, daß $\| {}_{(\alpha\beta)}\bar{b} \|$ (α, β Indizes des J_k Raumes) eine orthogonale Matrix ist. Aus (152), (153) folgt, daß

$$(155) \quad E_{p_1 p_2 \dots p_k} \mid q_1 q_2 \dots q_k {}_{(\beta)}\bar{A}^{q_1 q_2 \dots q_k} = {}_{(\beta)}\bar{A}_{p_1 p_2 \dots p_k}$$

ist, und aus (154'), (155) folgt dann, daß

$$\begin{aligned} (156) \quad {}_{(\beta\alpha)}\bar{b}_{(\gamma\alpha)}\bar{b} &= {}_{(\alpha)}A_{q_1 q_2 \dots q_k} {}_{(\alpha)}A_{p_1 p_2 \dots p_k} {}_{(\beta)}\bar{A}^{q_1 q_2 \dots q_k} {}_{(\gamma)}\bar{A}^{p_1 p_2 \dots p_k} = \\ &= E_{q_1 q_2 \dots q_k} \mid p_1 p_2 \dots p_k {}_{(\beta)}\bar{A}^{q_1 q_2 \dots q_k} {}_{(\gamma)}\bar{A}^{p_1 p_2 \dots p_k} = \\ &= {}_{(\beta)}\bar{A}^{q_1 q_2 \dots q_k} {}_{(\gamma)}\bar{A}_{q_1 q_2 \dots q_k} = \delta_{(\beta\gamma)}, \text{ w. z. b. w.} \end{aligned}$$

Es ist also ${}_{(\alpha)}A_{p_1 p_2 \dots p_k}$ bis auf orthogonale Transformationen (Äquivalenz) bestimmt; die Anzahl der linear unabhängigen Tensoren ${}_{(\alpha)}A_{p_1 p_2 \dots p_k}$ ist gleich dem Range der Matrix:

$$(157) \quad \| E_{p_1 p_2 \dots p_k} \mid q_1 q_2 \dots q_k \|,$$

also gleich l_k der Dimension des J_k . Die Tensoren (152') bestimmen bis auf Kongruenz die F_l im S_n .

Wir zeigen jetzt, daß man aus den Formen (149) $B_{pq} dy_p dy_q$, $B_{pqrs} dy_p dy_q dy_r dy_s, \dots$ die Tensoren (152') $E_{p|q}$, $E_{p_1 p_2 | q_1 q_2} \dots$ berechnen kann.

Wegen der Symmetrie $E_{p|q} = E_{q|p}$ folgt auf jedem Fall

$$(158) \quad E_{p|q} = B_{pq}.$$

Nehmen wir an, es wäre uns die Bestimmung der Tensoren (152') aus den Formen (149) bis zur $(2k-2)^{\text{ten}}$ Stufe gelungen (was auf jeden Fall für $k=2$ zutrifft), so können wir dann auch die Tensoren der $2k^{\text{ten}}$ Stufe $E_{p_1 p_2 \dots p_k} \mid q_1 q_2 \dots q_k$ berechnen. Aus (152'), (149) folgt nämlich:

$$(159) \quad B_{p_1 p_2 p_3 \dots p_k} \mid q_1 q_2 \dots q_k = \Sigma E_{\sigma_1 \sigma_2 \dots \sigma_k} \mid \sigma_{k+1} \dots \sigma_{2k},$$

wobei rechts über alle Permutationen $\sigma_1 \sigma_2 \dots \sigma_{2k}$ der Indizes $p_1 p_2 \dots p_k q_1 q_2 \dots q_k$ zu summieren ist. Nun gilt aber:

$$(160) \quad E_{p_1 p_2 \dots p_{k-1} p_k} | q_1 q_2 \dots q_{k-1} q_k - E_{p_1 p_2 \dots p_{k-1} q_k} | q_1 q_2 \dots q_{k-1} p_k = \\ = ({}^{(\alpha)}A_{p_1 p_2 \dots p_k} ({}^{(\alpha)}A_{q_1 q_2 \dots q_k} - ({}^{(\alpha)}A_{p_1 p_2 \dots p_k} {}^{(\alpha)}A_{q_1 q_2 \dots q_{k-1} p_k} {}^{34}),$$

d. h. die Differenz (160) ist der $(k-1)^{\text{te}}$ Krümmungstensor und ist also durch die Tensoren bis zur $(k-1)^{\text{ten}}$ Stufe ${}^{(\alpha)}A_p, {}^{(\alpha)}A_{pq} \dots {}^{(\alpha)}A_{p_1 p_2 \dots p_{k-1}}$ und deren Abteilungen bestimmt, sie sind also uns nach Voraussetzung bekannt. Aus (159) und (160) können wir demnach die Tensoren $E_{p_1 p_2 \dots p_k} | q_1 q_2 \dots q_k$ mit Hilfe der Tensoren $B_{pq} B_{pqrs}, \dots B_{p_1 p_2 \dots p_k} {}^{q_1 q_2 \dots q_k}$ berechnen (für jedes k) und aus den Tensoren $E_{p_1 p_2 \dots p_k} | q_1 q_2 \dots q_k$ die Formen ${}^{(\alpha)}A_{p_1 p_2 \dots p_k}$, wie wir gesehen haben.

Damit ist bewiesen, daß das System der Formenquadrate (149) die F_l bis auf ihre Lage im S_n fixiert, w. z. b. w.

Die Begriffe der h^{ten} Verbiegung³⁵⁾ und seine Anwendung auf die Hypertorse lassen sich mit Hilfe der Gleichungen ganz einfach auf den S_n Raum übertragen.

Zum Schluß wollen wir ganz kurz den Beweis des Satzes andeuten: Verschwindet für eine F_l im S_n die h^{te} Krümmung, so gibt es stets eine $\bar{F}_l$ im S_n , deren vollständiges Formensystem das System der ersten h Formen der F_l ist. (Über das vollständige Formensystem der F_l im R_n .)

Der Beweis wird geführt, wenn wir zeigen, daß das System:

$$(161) \quad \begin{cases} dx_i = (x)A_j ({}^{(\alpha)}\lambda^i dy_j \\ (d({}^{(\alpha)}\lambda^i)_R = ({}^{(\alpha\beta)}C_j ({}^{(\beta)}\lambda^i dy_j \end{cases}$$

(wobei α, β die Indizes der Räume J_k ($k=1, 2, \dots, h$) durchläuft), in welchen die Koeffizienten ${}^{(\alpha\beta)}C_j$ mit den entsprechenden der F_l übereinstimmen, total integrierbar ist.

Da die Integrabilitätsbedingungen der $J_1, J_2, \dots, J_{h-1}$ Räume des Systems (161) mit den entsprechenden für die F_l zusammenfallen, so sind sie identisch erfüllt. Es bleiben also nur noch zu untersuchen die Integrabilitätsbedingungen des letzten Raumes S_h (For-

³⁴⁾ Formenproblem, I. Kap., § 13, II. Kap., § 5.

³⁵⁾ Die Sätze I, II des II. Abschnitt, § 1, der Arbeit „Über das vollständige Formensystem der F_l im R_n “ in Monatsheften für Math. und Physik, Bd. XXX, lassen sich für die F_l im S_n auf dieselbe Weise beweisen. Berichtigung: Der Satz II (II. Abschnitt, § 1) der obengenannten Arbeit soll lauten: „Lassen sich zwei F_l so zuordnen, daß die entsprechenden Kurven gleiche Längen und gleiche Krümmungen bis zur $h-1^{\text{ten}}$ haben, so sind ihre h ersten Formen gleich, z. B. sind zwei F_2 im R_3 kongruent, wenn sie sich punktweise so zuordnen lassen, daß die entsprechenden Kurven gleiche Längen und gleiche erste Krümmung besitzen.“ Der Beweis dieses Satzes stützt sich wesentlich auf die Voraussetzung Formel [(13), (14)], daß entsprechende Kurven gleiche Längen besitzen.

menproblem, I. Kap., § 14). Es sind die sogenannten letzten Krümmungsgleichungen, die in unserem Falle mit den entsprechenden Integrabilitätsbedingungen der F_l übereinstimmen, da die h^{te} Krümmung der F_l verschwindet³⁶⁾. Es ist also das System (161) total integrabel, q. e. d.

Die $\overline{F}_e$ besitzt demnach die Einbettungszahl $k = \sum_{\sigma=1}^h l_\sigma$, wenn l_σ die Dimension des J_σ Raumes der F_l ist.

II. Nachtrag.

Es sei ein $V_n (n > 2)$, ein n -dimensionaler Riemannscher Raum gegeben, in dem es durch jeden Punkt in jeder Orientierung eine E_{n-1} gibt³⁷⁾. Seien dann v^i, w^i, λ^i drei linear unabhängige Vektoren — wobei $v_i \lambda^i = w_i \lambda^i = 0$ ist — in einem Punkt P , so gibt es eine E_{n-1} durch P mit λ^i als Normalvektor, die also die Vektoren v^i, w^i enthält.

Weiter liege in dieser E_{n-1} eine F_2 , welche in P das Flächenelement (v^i, w^i) enthält. In dieser F_2 denken wir uns eine geschlossene Kurve C von P nach P . Längs C führen wir λ^i parallel von P nach P . Da die Kurve C in der E_{n-1} , deren Normalvektor λ^i ist, liegt, so wird die Endlage von λ^i mit der Anfangslage von λ^i zusammenfallen³⁸⁾, so gilt

$$(162) \quad R^i{}_{jkl} \lambda_i v^l w^k = 0,$$

wobei

$$(162') \quad v_i \lambda^i = w_i \lambda^i = 0 \text{ ist.}$$

Der Tensor

$$(163) \quad R_{ijkl} v^k w^l = \mathcal{A}_{ij}(v, w)$$

läßt sich durch n linear unabhängige Vektoren ${}_{(\alpha)}\lambda^i$, wobei ${}_{(1)}\lambda^i = v^i$, ${}_{(2)}\lambda^i = w^i$, ${}_{(\beta)}\lambda^i {}_{(\gamma)}\lambda_i = \delta_{(\beta\gamma)}$ ($\beta, \gamma = 3, 4 \dots n$) und $v^i {}_{(\beta)}\lambda_i = w^i {}_{(\beta)}\lambda_i = 0$ ($\beta = 3, 4 \dots n$) ist, in der Form darstellen:

³⁶⁾ Formenproblem, I. Kap., § 14.

³⁷⁾ Die E_{n-1} enthält dauernd jeden Vektor der E_{n-1} , der aus einer in der E_{n-1} enthaltenen Anfangslage durch Parallelverschiebung längs beliebiger Kurven der E_{n-1} hervorgeht.

³⁸⁾ Aus der angegebenen Eigenschaft der E_{n-1} folgt: Die Parallelverschiebung des Vektors λ^i längs beliebiger Kurven der E_{n-1} ist eindeutig.

$$(164) \quad T_{ij}(v, w) = {}_{(\alpha\beta)}T(v, w) {}_{(\alpha)}\lambda_i {}_{(\beta)}\lambda_j$$

wobei

$$(164') \quad {}_{(\alpha\beta)}T(v, w) = - {}_{(\beta\alpha)}T(v, w) = {}_{(\alpha\beta)}A. \quad T_{ij}(v, w) {}_{(\alpha)}\lambda_i {}_{(\beta)}\lambda^i$$

ist³⁹⁾.

Es gilt dann zufolge (162), (162')

$$(165) \quad {}_{(\alpha\beta)}T(v, w) = R_{ijk} {}_{(1)}\lambda^k {}_{(2)}\lambda^l {}_{(\alpha)}\lambda^i {}_{(\beta)}\lambda^j = 0$$

für alle ${}_{(\alpha)}\lambda^i$, ${}_{(\beta)}\lambda^i$ sobald $\alpha \neq 1, 2$, oder $\beta \neq 1, 2$, ist, d. h.

$$(166) \quad {}_{(\alpha\beta)}T(v, w) = 0, \text{ für } \alpha \text{ oder } \beta > 2. \text{ Es ist}$$

also nur (166') ${}_{(12)}T(v, w) \neq 0$. Somit ist

$$(167) \quad T_{ij}(v, w) = \frac{1}{2} A(v, w) (v_i w_j - v_j w_i), \quad A(v, w) = {}_{(12)}T(v, w)$$

oder

$$(168) \quad R_{ijk} v^e w^k = \frac{1}{2} A(v, w) (g_i g_{jk} - g_{ik} g_j) v w^k.$$

Wir wollen jetzt zeigen, daß $A(v, w)$ von v^i, w^i unabhängig ist. Zu diesem Zwecke differenzieren wir (168) nach v^r und bekommen:

$$(169) \quad \begin{cases} R_{ijrk} w^k = \frac{1}{2} \frac{\partial A(v, w)}{\partial v^r} (g_i g_{jk} - g_{ik} g_j) v w^k + \\ + \frac{1}{2} A(v, w) (g_{ir} g_{jk} - g_{ik} g_{jr}) w^k \end{cases}$$

und multiplizieren (169) mit v^i, w^j :

$$(170) \quad \begin{cases} R_{ijrk} w^k v^i w^j = \frac{\partial A(v, w)}{\partial v^r} F(v, w) + \frac{1}{2} A(v, w) (g_{ir} g_{jk} - \\ - g_{ik} g_{jr}) w^k v^i w^j \end{cases} \quad (40)$$

Wegen $R_{ijrk} = R_{rkij}$ und wegen (168) gilt:

$$(171) \quad R_{ijrk} v^i w^j = \frac{1}{2} A(v, w) (g_{ri} g_{kj} - g_{rj} g_{ki}) v^i w^j$$

³⁹⁾ ${}_{(\alpha\beta)}T(u, v) = [T_{ij}(u, v) {}_{(\alpha)}\lambda^i {}_{(\beta)}\lambda^j] : ({}_{(\alpha)}\lambda^i {}_{(\alpha)}\lambda_i) \cdot ({}_{(\beta)}\lambda^k {}_{(\beta)}\lambda_k),$
 ${}_{(\alpha\alpha)}T(u, v) = 0.$

⁴⁰⁾ wobei $F(u, v) \neq 0$ ist, da v^i, w^i voneinander linear unabhängig sind.

Aus (171) und (170) folgt demnach:

$$(172) \quad \frac{\partial A(v, w)}{\partial v^r} = 0 \text{ d. h. } A(v, w) \text{ ist von } v^i \text{ unabhängig (auf die}$$

selbe Weise zeigt man, daß $A(v, w)$ auch von w^i unabhängig ist). Da aber v^i, w^i willkürliche Vektoren sind, so folgt dann aus (171):

$$(173) \quad R_{i j k e} = A(g_{i e} g_{j k} - g_{i k} g_{j e}).$$

Aus der Bianchischen Identität folgt dann, daß A von $x_1, x_2, \dots, x_n$ unabhängig ist, also eine Konstante ist. Es ist demnach V_n ein Raum von konstanter Krümmung.

(Eingegangen: 24. V. 1928.)

Die Starrheit der Eiflächen.

Von Adalbert Duschek in Wien.

Der Satz, daß eine Eifläche, also eine geschlossene Fläche mit überall positiver und stetiger Gauss'scher Krümmung, als Ganzes nicht verbogen werden kann, wurde bekanntlich zum erstenmal von Liebmann¹⁾ aufgestellt, dessen Beweis dann später von Weyl²⁾ und Blaschke³⁾ wesentlich vereinfacht und verbessert wurde. Ich habe vor einigen Jahren darauf hingewiesen⁴⁾, daß sich der Satz recht einfach und kurz mit Hilfe einiger Formeln für die Relativkrümmungen einer Fläche bezüglich einer zweiten Fläche beweisen läßt. Es bietet nun keine Schwierigkeit, diesen Beweis ganz aus den Überlegungen der relativen Flächentheorie herauszulösen; er ist in dieser unabhängigen Form noch immer kürzer und, wie ich glaube, auch durchsichtiger als die früheren. Der wesentlichste Unterschied gegenüber Blaschkes Beweis besteht in der Deutung des Vektors η als Ortsvektor einer zweiten Fläche und in der Heranziehung der Formel von Gauss-Bonnet.

Es sei $\mathfrak{x}(u, v)$ die vektorielle, dreimal stetig differenzierbare Parameterdarstellung einer Eifläche. Wir nehmen an, die Fläche ließe sich verbiegen und denken sie uns als Fläche $t = 0$ in der Schar

$$(1) \quad \mathfrak{x} = \mathfrak{x}(u, v, t)$$

der aus ihr durch Verbiegungen entstehenden Flächen eingebettet. (1) sei nach t mindestens einmal stetig differenzierbar, man kann dann für hinreichend kleine t

$$(2) \quad \mathfrak{x}(u, v, t) = \mathfrak{x}(u, v, 0) + t \frac{\partial \mathfrak{x}}{\partial t}(u, v, \delta t), \quad (0 < \delta < 1)$$

schreiben. Zur Abkürzung sei

$$(3) \quad \mathfrak{x}(u, v, t) = \overline{\mathfrak{x}}, \quad \mathfrak{x}(u, v, 0) = \mathfrak{x}, \quad \frac{\partial \mathfrak{x}}{\partial t}(u, v, 0) = \mathfrak{z}$$

gesetzt. Dann ist für sehr kleine t in erster Annäherung

$$(4) \quad \overline{\mathfrak{x}} = \mathfrak{x} + t \mathfrak{z}$$

und

$$(5) \quad d\overline{\mathfrak{x}} = d\mathfrak{x} + t d\mathfrak{z},$$

¹⁾ Göttinger Nachrichten 1899; Math. Annalen 53 (1900) und 54 (1901).

²⁾ Berliner Sitzungsberichte 1917.

³⁾ Göttinger Nachrichten 1912, Math. Zeitschrift 9 (1921), Differentialgeometrie I, Seite 131 f.

⁴⁾ Wiener Sitzungsberichte IIa, 135 (1926).

so daß sich als Maß für die Änderung des quadrierten Bogenelementes beim Übergang von der Fläche $\mathfrak{x}$ zur Fläche $\overline{\mathfrak{x}}$

$$(6) \quad \frac{d(ds^2)}{dt} = \lim_{t \rightarrow 0} \frac{d\overline{\mathfrak{x}}^2 - d\mathfrak{x}^2}{t} = 2 d\mathfrak{x} d\mathfrak{z}$$

ergibt; soll also $\overline{\mathfrak{x}}$ aus $\mathfrak{x}$ durch Verbiegung entstehen, so muß

$$(7) \quad d\mathfrak{x} d\mathfrak{z} = 0$$

sein, und zwar identisch in $du:dv$, d. h. für alle Richtungen.

Durch (7) ist in jedem Flächenpunkt (u, v) ein zweimal stetig differenzierbarer, von $du:dv$ unabhängiger Vektor $\mathfrak{y}$ bestimmt, für den

$$(8) \quad d\mathfrak{z} = \mathfrak{y} \times d\mathfrak{x}$$

und somit auch (7) identisch in $du:dv$ gilt; aus (5) folgt dann

$$d\overline{\mathfrak{x}} - d\mathfrak{x} = t \mathfrak{y} \times d\mathfrak{x}$$

und somit ist $t\mathfrak{y}$ nichts anderes als der Drehvektor für die infinitesimale Verbiegung der Fläche $\mathfrak{x}$ in die Fläche $\overline{\mathfrak{x}}$, d. h. Vektor, dessen Richtung die Richtung der Achse jener Drehung angibt, die das Flächenelement im Punkt (u, v) bei Verbiegung erfährt, während Länge und Sinn des Vektors den Drehwinkel bis auf Größen höherer Ordnung festlegen⁵⁾. — Ist unser Satz richtig, so muß $\mathfrak{y}$ von u und v unabhängig sein, d. h. die Fläche sich höchstens als starrer Körper bewegen, aber nicht verbiegen lassen. Unser Beweis läuft auch darauf hinaus, zu zeigen, daß die Annahme, $\mathfrak{y}$ sei nicht konstant, auf einen Widerspruch führt.

Aus (8) folgt

$$(9) \quad \mathfrak{z}_u = \mathfrak{y} \times \mathfrak{x}_u, \quad \mathfrak{z}_v = \mathfrak{y} \times \mathfrak{x}_v;$$

differenziert man die erste Gleichung nach v , die zweite nach u und subtrahiert, so erhält man die Integrabilitätsbedingung

⁵⁾ Ist $\mathbf{e}$ ein Einheitsvektor in der Drehachse, α der Drehwinkel in positivem Sinn und $\mathbf{p}$ ein beliebiger Vektor mit dem Anfangspunkt auf der Drehachse, so geht $\mathbf{p}$ bei der Drehung über in

$$\mathbf{p} = \mathbf{p} \cos \alpha + \mathbf{e} (\mathbf{e} \cdot \mathbf{p}) (1 - \cos \alpha) + \mathbf{e} \times \mathbf{p} \sin \alpha.$$

Ist α sehr klein, so wird in erster Annäherung (infinitesimale Drehung)

$$\overline{\mathbf{p}} = \mathbf{p} + \mathbf{e} \times \mathbf{p} \alpha.$$

Als Drehvektor bezeichnet man dann in der Regel den infinitesimalen Vektor $\alpha \mathbf{e}$. Der Verbiegungsparameter t des Textes ist jedenfalls von derselben Größenordnung wie der Drehwinkel.

$$(10) \quad \mathfrak{x}_u \times \mathfrak{y}_v = \mathfrak{x}_v \times \mathfrak{y}_u$$

aus der man durch skalare Multiplikation mit $\mathfrak{x}_u$ und $\mathfrak{x}_v$ leicht folgert, daß die vier Vektoren $\mathfrak{x}_u, \mathfrak{x}_v, \mathfrak{y}_u, \mathfrak{y}_v$ in einer Ebene liegen, weshalb (10) nur eine einzige Bedingung darstellt.

Deutet man also $\mathfrak{y}(u, v)$ als Ortsvektor einer zweiten Fläche, so ist diese auf $\mathfrak{x}(u, v)$ umkehrbar eindeutig durch parallele Normalen abgebildet. Somit existieren vier stetig differenzierbare Funktionen a, b, c, d , so daß

$$(11) \quad \mathfrak{y}_u = a \mathfrak{x}_u + b \mathfrak{x}_v, \quad \mathfrak{y}_v = c \mathfrak{x}_u + d \mathfrak{x}_v$$

ist. Die Bedingung (10) gibt

$$(12) \quad a + d = 0.$$

Die Abbildung von $\mathfrak{y}$ auf $\mathfrak{x}$ ist gegensinnig (d. h. der Umlaufsinn entsprechender geschlossener Kurven auf den beiden Flächen ist entgegengesetzt) oder, was auf dasselbe herauskommt, die Gaußschen Krümmungen der beiden Flächen haben in entsprechenden Punkten entgegengesetztes Vorzeichen. Zum Beweis dieser Behauptung denken wir uns die Fläche $\mathfrak{x}$ auf ihre Krümmungslinien als Parameterkurven bezogen (unter Beibehaltung der bisherigen Bezeichnung). Differenziert man (11) nach u und v , so erhält man durch skalare Multiplikation mit der gemeinsamen Flächennormalen Relationen zwischen den Koeffizienten L, M, N und L', M', N' der beiden zweiten Grundformen von $\mathfrak{x}$ und $\mathfrak{y}$, nämlich — noch in allgemeinen Parametern —

$$(13) \quad \begin{aligned} L' &= a L + b M \\ M' &= a M + b N = c L + d M \\ N' &= c M + d N \end{aligned}$$

woraus insbesondere

$$(14) \quad c L + (d - a) M - b N = 0$$

folgt. Für die Krümmungslinien als Parameterkurven wird $M = F = 0$ (E, F, G sind wie gewöhnlich die Koeffizienten der ersten Grundform von $\mathfrak{x}$). Die Krümmung K von $\mathfrak{x}$ wird somit

$$(15) \quad K = \frac{L N}{E G},$$

wegen $K > 0$ ist also auch $LN > 0$. Aus (14) folgt

$$c L - b N = 0,$$

also ist $bc \geq 0$. (13) geben

$$(16) \quad L' = aL, \quad M' = cL = bN, \quad N' = dN,$$

also

$$(17) \quad L'N' - M'^2 = (ad - bc)LN.$$

Ist K' die Gauss'sche Krümmung von η , so folgt weiter

$$(18) \quad \text{sign } K' = \text{sign } (L'N' - M'^2)$$

oder wegen (17) und (12)

$$(19) \quad \text{sign } K' = \text{sign } (-a^2 - bc) = -1,$$

da a jedenfalls reell und, wie schon gezeigt, $bc \geq 0$ ist.

Da $K > 0$ ist, muß $K' < 0$ sein, d. h. die Fläche η ist in allen Punkten negativ gekrümmt; wegen der ein-eindeutigen Abbildung auf ξ ist sie ferner eine (geschlossene) Fläche von denselben Zusammenhangsverhältnissen wie ξ , also vom Zusammenhang der Kugel. Nach der Gauss-Bonnetschen Formel ist das über die ganze Fläche η erstreckte Integral der Gauss'schen Krümmung

$$(20) \quad \int K' do = 4\pi$$

also jedenfalls positiv, in Widerspruch dazu, daß bei $K' < 0$ auch $\int K' do < 0$ ist.

Wie schon eben angekündigt, bleibt nur der Schluß, daß η konstant ist. Aus (8) erhält man durch Integration

$$(21) \quad \mathfrak{z} = \mathfrak{z}_0 + \eta \times \xi,$$

so daß (4) zu

$$(22) \quad \bar{\xi} = \xi + t(\mathfrak{z}_0 + \eta \times \xi)$$

wird. Diese Formel gilt natürlich nur näherungsweise für sehr kleine t , aber, wie man leicht nachrechnet, ändert sich der Abstand zweier beliebiger Punkte der Fläche ξ bei der Transformation (22) nur um Größen, die in t von mindestens zweiter Ordnung sind, d. h. aber, daß die Fläche wie ein starrer Körper bewegt wird.

(Eingegangen: 28. X. 1928.)

On continua which are disconnected by the omission of any point and some related problems¹⁾.

By W. L. Ayres²⁾ in Vienna.

1. Continua which are disconnected by the omission of any point have already been studied by J. R. Kline³⁾. He has shown (1) that any such continuum is an unbounded, acyclic (i. e. containing no simple closed curve) continuous curve, (2) that each point of the continuum is contained in an open curve of the continuum, and further (3) that within any bounded region there are only a finite number of cut points of power greater than two⁴⁾.

Let us take the second of these properties for consideration. A continuous curve having this property may be thought of as a web stretched between certain points at infinity and, accordingly, we shall say a set is a web if it is a continuous curve such that each point of the curve belongs to an open curve which is a subset of the curve. If a web contains no simple closed curve we shall say that it is an acyclic web. It follows evidently that every point of an acyclic web is a cut point, and if every point of a continuum is a cut point, then it is an acyclic web. Hence this is an optional definition of the acyclic webs.

In the first section of this paper we shall show that there exists in the plane a universal acyclic web, that is, an acyclic web such that every acyclic web is homeomorphic with some subset of this plane acyclic web. In the second section we shall derive a property of a point if it is a non-cut point of the continuous curve. In the third we shall find some properties which characterize the acyclic webs. In the last section we shall consider some webs which are subsets of general continuous curves.

¹⁾ Presented to the American Mathematical Society April 7, 1928.

²⁾ National Research Fellow in Mathematics.

³⁾ Closed connected sets which are disconnected by the removal of a finite number of points, *Proceedings of the National Academy*, vol. 9 (1923), pp. 7—12.

⁴⁾ If P is a point of a connected set M , P is said to be a cut point of power n if $M - P$ consists of exactly n components.

Unless stated to the contrary, we assume every point set considered in this paper to lie in a Euclidean space E_n of n dimensions.

Notation. We shall use the ordinary notation of the theory of sets, such as $A+B$, $A-B$, $A \cdot B$ etc. If K is a set, $\rho(K)$ denotes the diameter of K . If C is a hypersphere, i. e. the set of points whose distance from some point P is a constant r , then $I(C)$ and $E(C)$ denote the interior and exterior of C , i. e. the set of points whose distance from P is less than r and greater than r respectively.

I. A Universal Acyclic Web.

2. Theorem 1. There exists in the plane an acyclic web W such that every acyclic web is homeomorphic with some subset of W .

In a plane E_2 , let K denote the positive half of the x -axis. Let K_{1ij} ($j = \pm 1$) denote the point set consisting of the segment from $(1,0)$ to $(2,j)$ together with the points of the line $y=j$ for which $x \geq 2$. Let K_{2ij} ($i = 0, \pm 1; j = \pm 1, \pm 2$) denote the set consisting of the line segment from $(3,i)$ to $(4, i + [j/5])$ together with all points of the line $y = i + (j/5)$ for which $x \geq 4$. Let K_{3ij} ($i = 0, \pm 1, \pm 2, \dots, \pm 7; j = \pm 1, \pm 2, \pm 3$) denote the set consisting of the line segment from $(5, i/5)$ to $(6, [i/5] + [j/5 \cdot 7])$ together with all points of the line $y = i/5 + j/5 \cdot 7$ for which $x \geq 6$. Continue this process. In general let K_{nij} ($i = 0, \pm 1, \pm 2, \dots, \pm [(2n-1)R_{n-1} + n-1]; j = \pm 1, \pm 2, \dots, \pm n$) denote the set consisting of the straight line segment from $(2n-1, i/\pi_n)$ to $(2n, i/\pi_n + j/\pi_{n+1})$ together with the points of the line $y = i/\pi_n + j/\pi_{n+1}$ for which $x \geq 2n$, where π_k ($k = n, n+1$) denotes the product $5 \cdot 7 \cdot 9 \dots (2k-1)$ and R_{n-1} denotes the largest value of i for which the set $K_{n-1, i, j}$ is defined. Let

$$M = K + \sum_{n=1}^{\infty} \sum_{i=-R_n}^{+R_n} \sum_{j=-n}^{+n} K_{nij} {}^5).$$

Let W be the set of points consisting of M together with the set obtained from M by reflection in the y -axis. We shall show that the set W has the desired properties.

It is evident that the set W is an acyclic web. Let U be any acyclic web. We will complete our proof by showing that there exists an acyclic web V and a (1,1) correspondence T such that V is a subset of W , $T(U) = V$, and both T and T^{-1} are continuous. Let L_1 be any open curve in U and let N_1 be the x -axis in the plane E_2 . In any finite region there are on L_1 only a finite number of points of U of power greater than two. Let P_0 be a point of

⁵⁾ Σ' denotes that the summation is over all values included between the limits except zero.

L_1 whose power in U is two, and let Q_0 be the origin in the plane E_2 . From property (3) proved by Kline we may show that the points of U which are cut points of power greater than two and lie on L_1 may be designated by the letters $\dots P_{-3}, P_{-2}, P_{-1}, P_1, P_2, P_3, \dots$ in such a manner that P_j is between P_i and P_k if, and only if, $i < j < k$ or $k < j < i$, and further P_0 is between P_i and P_{-i} for each i . On the positive half of the x -axis let Q_1 be the point of smallest abscissa such that Q_1 is a cut point of W of power greater than the power of P_1 in U . Let Q_2 be the point of smallest abscissa such that Q_1 is between Q_0 and Q_2 and the power of Q_2 in the set W is greater than the power of the point P_2 in the set U . Let us continue to define a point Q_n for each P_n in this manner. In exactly a similar manner we may define a point Q_{-n} on the negative half of the x -axis for each P_{-n} . For these two sets we define T such that $T(P_k) = Q_k$ ($k = \dots, -2, -1, 0, 1, 2, \dots$). Set up any correspondence between the arc $P_k P_{k+1}$ of the open curve L_1 and the straight-line interval $Q_k Q_{k+1}$ of N_1 , which is one-to-one and continuous in both directions and such that P_k and Q_k correspond and P_{k+1} and Q_{k+1} correspond; and we shall call this correspondence T for these two arcs. If there are only a finite number of the points P_n , let P_s be the one of largest subscript. Then set up any continuous (1,1) correspondence between that ray⁶⁾ of L_1 which has P_s as its vertex and contains no other point P_k and that ray of N_1 which has Q_s as its vertex and contains no other point Q_k , and call this correspondence T for these two rays. We set up a similar correspondence between rays if there are only a finite number of the points P_{-n} . In this way we have defined a continuous (1,1) correspondence between the entire open curves L_1 and N_1 such that for any k for which P_k is defined, we have $T(P_k) = Q_k$. Now suppose there are m_k components, $C_{k1}, C_{k2}, C_{k3}, \dots, C_{km_k}$, of $U - L_1$ which have P_k as a limit point. Each of the sets $C_{ki} + P_k$ contains a ray L_{ki} with vertex P_k ⁷⁾. Since Q_k is a cut point of W of power greater than that of P_k in U , the set $K_{n,i}, -n, K_{ni}, -n+1, \dots, K_{ni}, -1, K_{ni}, 1, \dots, K_{ni}, n$, when n and i have the value such that $(2n-1, i/\pi_n)$ is the point Q_k , contains more than m_k rays. Let $N_{k1}, N_{k2}, \dots, N_{km_k}$ be any m_k of them. On the ray L_{ki} ($1 \leq i \leq m_k$) let $P_{ki1}, P_{ki2}, \dots$ be the points which are cut points of U of power greater than two and such that we have the order $P_k P_{ki1} P_{ki2} \dots$ on the ray L_{ki} . In the same way as above there is on N_{ki} a sequence of points $Q_{ki1}, Q_{ki2}, Q_{ki3}, \dots$, where

⁶⁾ If X is any point of an open curve C , the set $C - X$ consists of two components or maximal connected subsets. Each of these components plus the point X is called a ray and X is said to be the vertex of each ray. Or a ray may be defined as a closed point set in continuous (1,1) correspondence with those points of the x -axis for which x is non-negative.

⁷⁾ B. Knaster and C. Kuratowski, Sur les continus non-bornés. Fund. Math., vol. 5 (1924), p. 26.

Q_{ki_1} is defined as the point of the set $N_{ki} - Q_k$ of smallest abscissa such that the power of Q_{ki_1} in W is greater than the power of P_{ki_1} in U , and Q_{kij} is defined as the point of N_{ki} of smallest abscissa such that Q_{kij-1} is between Q_k and Q_{kij} and the power of Q_{kij} in W is greater than the power of P_{kij} in U . As above we may define a continuous (1,1) correspondence T between L_{ki} and N_{ki} such that

$$T(P_k) = Q_k \text{ and } T(P_{kij}) = Q_{kij}$$

for each j for which P_{kij} is defined. We repeat this process exactly as above by considering the components $[L_{kijm}]$ of the set $U - L_1 - \sum_k \sum_i L_{ki}$ that have P_{kij} as a limit point. It is clear how the process is continued from this point and without difficulty we may see that in this way we obtain an acyclic web

$$V = N_1 + \sum_k \sum_i N_{ki} + \sum_k \sum_i \sum_j \sum_m N_{kijm} + \dots,$$

which is a subset of the acyclic web W and a continuous (1,1) correspondence T such that $T(U) = V$.

3. Remarks. We may define an acyclic web to be of power n (or order n , since for acyclic webs the power and the order of a point is the same) if it contains one cut point of power n but no cut point of power greater than n . To obtain a universal acyclic web of power n , we may follow the construction of W as given above except that in the stages of the construction where so many rays are added that a cut point of power greater than n is produced, we may add only enough rays to obtain a cut point of power n .

If each arc between two points of W of power greater than two is replaced by a Wazewski-Menger universal curve⁸), that is, the plane bounded acyclic continuous curve as defined by Wazewski and Menger such that every bounded acyclic continuous curve is homeomorphic with some subset of it, and at each point of W of power greater than two there is attached a countable infinity of these Wazewski-Menger curves, each having just one point in common with W and no two having any other point in common and such that their diameters converge toward zero; then with this example we may establish the following result: There exists in the plane an acyclic continuous curve such that every acyclic continuous curve (bounded or not) is homeomorphic with some subcontinuum of this plane curve.

⁸) T. Wazewski, Sur les courbes de Jordan ne renfermant aucune courbe simple fermée de Jordan. Ann. Soc. Pol. Math., vol. 2 (1923), pp. 49—170; K. Menger, Zur allgemeinen Kurventheorie. Fund. Math., vol. 10 (1926), p. 108.

II. The Non-cut Points of a Continuous Curve.

In studying the acyclic webs it will be useful to know some properties of a continuous curve if a point fails to cut it. Accordingly we shall prove the following property:

4. **Theorem 2⁹**. If P is a non-cut point of a continuous curve M (in space of n dimensions) and ε is any positive number, there exists a continuous curve N_ε and a positive number δ_ε such that

$$(a) N_\varepsilon \subset M,$$

$$(b) \rho(N_\varepsilon) < \varepsilon,$$

(c) every point of M whose distance from P is less than δ_ε belongs to the set N_ε ,

(d) the set $M - N_\varepsilon$ is connected.

Proof. Let C_1 and C_2 be hyperspheres with center P and radii $\varepsilon/4$ and $\varepsilon/5$ respectively. It is easy to see that there are but a finite number of the components of $M - M \cdot (C_2 + I[C_2])$ that contain points of $E(C_1)$. Let these be denoted by $M_1, M_2, \dots, M_s$. For each value of i ($1 \leq i \leq s$) let P_i be a point of M_i . Since P is a non-cut point of M , for each value of i ($1 \leq i < s$) there exists an arc α_i lying wholly in $M - P$ and with end points P_i and P_{i+1} ¹⁰. There exists a hypersphere C_3 with center P and radius $r < \varepsilon/5$ such that $E(C_3)$ contains $\Sigma \alpha_i$. Let C_4 be a hypersphere with center P and radius $\frac{1}{2}r$. Let K denote the component of $M \cdot (C_4 + I[C_4])$ containing P . By a theorem due to G. T. Whyburn and the author¹¹, there exists a continuous curve H which contains K and is a subset of M and such that every point of H is within a distance $r/4$ of some point of K . Hence

$$H \subset I(C_3) \text{ and } H \cdot \sum_{i=1}^{s-1} \alpha_i = 0.$$

⁹) H. M. Gehman has recently announced a result which is very similar to Theorem 2. If the phrase „continuous curve N_ε “ in the statement of Theorem 2 is replaced by „ M -domain N_ε “, we have exactly Gehman's result. A connected subset K of a continuum M is said to be an M -domain if $M - K$ is closed. These two results, while very much alike, were discovered independently; and, in fact, were announced to the American Mathematical Society on the same day. For abstracts of these two papers, see Bull. Amer. Math. Soc., vol. 34 (1928), pp. 428 and 430. With the use of this theorem of Gehman's, the proofs of parts (1) and (4) of Theorem 3 may be shortened slightly.

¹⁰) R. L. Moore, Concerning continuous curves in the plane, Math. Zeit., vol. 15 (1922), p. 255. Although stated for two dimensions, the proof of this theorem holds equally well in n dimensions.

¹¹) On continuous curves in n dimensions, Bull. Amer. Math. Soc., vol. 34 (1928), pp. 349–360, theorem 1.

Now let

$$N_\varepsilon = H + \Sigma D,$$

where ΣD denotes the sum of all those components of $M-H$ that contain no points of $E(C_1)$. By a theorem due to H. M. Gehman¹²⁾ the set N_ε is a continuous curve. By definition $I(C_1) + C_1$ contains N_ε and thus $\rho(N_\varepsilon) \leq \frac{1}{2}\varepsilon < \varepsilon$. Since M is connected im kleinen at P there exists a positive number δ_ε such that every point of M whose distance from P is less than δ_ε lies with P in a connected subset of M every point of which is within a distance $\frac{1}{2}\varepsilon$ of P . Hence every point of M whose distance from P is less than δ_ε belongs to K and thus to N_ε .

It remains to see that $M-N_\varepsilon$ is connected. The set $M-N_\varepsilon$ consists of those components of $M-H$ that contain points of $E(C_1)$. Since $C_2 + I(C_2)$ contains H , each such component contains one of the sets M_i . Further

$$M-N_\varepsilon \supset \sum_{i=1}^{s-1} \alpha_i.$$

Then since $\Sigma M_i + \Sigma \alpha_i$ is a connected subset of $M-N_\varepsilon$ and contains a point of each component of $M-N_\varepsilon$, the set $M-N_\varepsilon$ is connected.

Corollary. If P is a cut point of the continuous curve M of power n and ε is any positive number, there exists a positive number δ_ε and a continuous curve N_ε such that conditions (a) to (c) of Theorem 2 are satisfied and (d) $M-N_\varepsilon$ consists of exactly n components.

5. The property given in theorem 2 does not hold for general continua even if N_ε is required to be merely a continuum in place of a continuous curve. Consider the following example: In a plane E_2 with polar coordinates (ρ, Θ) , let K_1, K_2, K_3, K_4 be the intervals of the lines $\Theta = 0, \Theta = \pi/2, \Theta = \pi, \Theta = 3\pi/2$ from $\rho = 0$ to $\rho = 1$, and for each $n > 4$ let K_n be the interval of the line $\Theta = 2\pi(1 - 1/2^{n-2})$ from $\rho = 0$ to $\rho = 1$. On each interval K_i ($i = 1, 2, 3, \dots$) let P_{ij} be the point such that $\rho = 1/j$ ($j = 2, 3, \dots$). Let K_{ij} be the interval of the line through P_{ij} parallel to K_{i+1} from $\rho = 1/j$ to $\rho = 1$. Let

$$M = \sum_{i=1}^{\infty} (K_i + \sum_{j=2}^{\infty} K_{ij}).$$

Then M is a continuum and the point Q ($\rho = 0$) is a non-cut point of M . But if ε is any positive number less than one, and N

¹²⁾ Some relations between a continuous curve and its subsets, *Annals of Mathematics*, Ser. 2, vol. 28 (1927), pp. 103-111, theorem 8.

is any subcontinuum of M whatsoever, which contains Q and is of diameter less than ε , then $M-N$ is not connected.

In this example we see that the property of non-cut points given in theorem 2 does not hold true in general continua even when condition (c) is omitted. This property of the non-cut points may be used to characterize the continuous curves even with condition (d) omitted. Conditions (a) to (c) simply make M connected im kleinen at the point P . Stated in this form, this characterization is that a continuum M is a continuous curve if, and only if, it is connected im kleinen at each of its non-cut points.

III. The Acyclic Webs.

6. In this section we will characterize the acyclic webs through several properties. In view of the result of section I of this paper, there is no loss of generality in confining ourselves to two dimensions for the study of acyclic webs. Accordingly it will be understood that the point sets mentioned in this section lie in a Euclidean two-space E_2 .

7. **Theorem 3.** In order that a continuous curve M be an acyclic web any of the following conditions are necessary and sufficient:

- (1) if K is any proper subcontinuum of M , then $M-K$ be unbounded;
- (2) if D is any proper M -domain of M , then $M-D$ be unbounded;
- (3) every bounded subcontinuum of M cut M ;
- (4) every bounded M -domain of M cut M ;
- (5) (a) M be non-dense, (b) the boundary of every complementary domain of M be an open curve.

Proof. The proof that each of the conditions is necessary affords no difficulty. We shall now show that each gives a sufficient condition.

(1) Suppose M contains a point P which is not a cut point of M . There exists a continuous curve N_ε and a positive number δ_ε satisfying conditions (a) to (d) of theorem 2, where ε is any positive number. The set $M-N_\varepsilon$ is connected and let $K=\overline{M-N_\varepsilon}$. The set K is a subcontinuum of M and, from condition (c) of theorem 2, the set $M-K$ is non-vacuous. The set $M-K$ is a subset of N_ε and thus is bounded since $\rho(N_\varepsilon) < \varepsilon$. But this is contrary to the assumed condition.

(2) Suppose M contains a point which is not a cut point of M . There exists a continuous curve N_ε and a positive number δ_ε satis-

fying (a) to (d) of theorem 2, where ε is any positive number. Since N_ε is closed and $M - N_\varepsilon$ is connected, the set $M - N_\varepsilon$ is an M -domain. But $M - (M - N_\varepsilon) = N_\varepsilon$ is bounded, contrary to our condition (2).

(3) It follows by Theorem 2 that if M contains a non-cut point, there exists a bounded subcontinuum N_ε of M such that $M - N_\varepsilon$ is connected, contrary to our condition.

(4) In (1) we showed that the assumption that M contains a non-cut point P leads to the existence of a continuum K such that $M - K$ contains P and is bounded. Let L be the component of $M - K$ containing P . By a theorem due to Kuratowski and Hahn¹³ the set L is an open subset of M and thus L is a bounded M -domain. It is easy to show that every component of $(M - K) - L$ has a limit point in the set K . Hence $M - L$ is connected and this contradicts condition (4).

(5) Suppose M contains a simple closed curve J . If the interior of J belongs to M , then condition (a) is contradicted. If there is in the interior of J a point P which belongs to $E_2 - M$, let B be the boundary of the complementary domain of M containing P . Since M contains J , the set B is a subset of J plus its interior and hence cannot be an open curve. Therefore M cannot contain any simple closed curve. Every point of an acyclic continuous curve is a cut point, except for the end points. Suppose Q is an end point of M . There exists a simple closed curve J_1 enclosing the point Q and having just one point X in common with M ¹⁴. Let D be the complementary domain of M containing $J_1 - X$ and let B_1 denote its boundary. By condition (b) the set B_1 is an open curve. Hence every component of $B_1 - X$ is unbounded. But there is at least one component of $B_1 - X$ which lies in the interior of J_1 and thus is bounded. Hence M has no end points and every point of M is a cut point of M .

8. **Theorem 4.** In order that a continuum M be an acyclic web either of the following conditions is necessary and sufficient:

(1) Every point of M be accessible from two unbounded complementary domains of M ;

(2) (a) M be non-dense, (b) the boundary of every complementary domain of M be an open curve, (c) if J is any simple closed curve there are only a finite number of complementary domains of M which have points interior to J .

¹³ C. Kuratowski, Une définition topologique de la ligne de Jordan, *Fund. Math.*, vol. 1 (1920), pp. 40-43; H. Hahn, Über die Komponenten offener Mengen, *Fund. Math.*, vol. 2 (1921), pp. 189-192.

¹⁴ H. M. Gehman, Concerning end points of continuous curves and other continua, *Trans. Amer. Math. Soc.*, vol. 30 (1928), pp. 63-84, theorem 2.

Proof. (1) With the use of a theorem due to R. G. Lubben¹⁵⁾, the condition may be shown to be necessary. We shall show that the condition is also sufficient. Since every point of M is accessible from two complementary domains, M is a continuous curve by a result due to G. T. Whyburn¹⁶⁾. Now let P be any point of M and let D_1 and D_2 be two unbounded complementary domains of M from which P is accessible. Since M is a continuous curve, it may be shown that there exist two rays R_1 and R_2 each of which has P as its vertex and such that R_i lies in D_i except for P . Let D' and D'' be the two complementary domains of $R_1 + R_2$. Since D' contains points of both D_1 and D_2 , it must contain a point of M for otherwise the domains D_1 and D_2 are identical. Similarly D'' contains a point of M . Therefore P is a cut point of M for

$$M - P = D' \cdot M + D'' \cdot M,$$

and $D' \cdot M$ and $D'' \cdot M$ are two non-vacuous mutually separated sets.

(2) The necessity of the conditions is evident. We shall show that they are sufficient. If the boundary of every complementary domain is an open curve, then every boundary point of a complementary domain is accessible from all sides from the domain. With this and condition (c) we have the conditions for the Schoenflies characterization of a generalized continuous curve¹⁷⁾. Hence M is a continuous curve. Using conditions (a) and (b) we may show now that every point of M is a cut point just as in the fifth part of Theorem 3.

9. Remarks. It is interesting to notice that while every point of M is a cut point if every point of M is accessible from two unbounded complementary domains, nevertheless it is not true that if a single point P is accessible from two unbounded complementary domains of M then P must be a cut point of M . Consider the following example: Let M consist of those points of the y -axis

¹⁵⁾ The separation of mutually separated subsets of a continuum by curves, Bull. Amer. Math. Soc., vol. 32 (1926), p. 114 (abstract).

¹⁶⁾ Concerning accessibility in the plane and regular accessibility in n dimensions, Bull. Amer. Math. Soc., vol. 34 (1928), pp. 504—510, theorem 2.

¹⁷⁾ A. Schoenflies, Die Entwicklung der Lehre von den Punktmannigfaltigkeiten, zweiter Teil, Jahr. der Deutschen Math.-Ver., Ergänzungsbände, vol. 2 (1908), p. 237. The generalized form of the Schoenflies theorem is: A plane continuum M (bounded or not) is a continuous curve if, and only if, (1) every point of M which lies on the boundary of some complementary domain of M is accessible from all sides from the domain, (2) if ε is any positive number and J is any simple closed curve then only a finite number of the complementary domains of M of diameter greater than ε have points within J .

for which $y \geq -1$ together with the intervals from $(1/\pi, -1)$ to $(1/\pi, 0)$ and $(0, -1)$ and the points of the curve

$$y = \frac{1}{x} \sin^2 \frac{1}{x} \quad \text{for } 0 < x \leq \frac{1}{\pi}.$$

In this example all the points of the interval from $(0, -1)$ to $(1/\pi, -1)$ are accessible from two unbounded complementary domains and yet not one of these points is a cut point of M .

Consider the following example: For each positive integer n let H_n denote the interval from $(1/n, 1/n)$ to $(1/n, -1/n)$, K_n denote the set of all points (x, y) for which $x < 1/n$, $y = \pm 1/n$, N_n denote the interval from $(-n, 1/n)$ to $(-n, 1/[n+1])$. Let K_0 denote the set of all points (x, y) for which $x \leq 0$, $y = 0$. Let

$$M = \sum_{n=0}^{\infty} K_n + \sum_{n=1}^{\infty} H_n + \sum_{n=1}^{\infty} N_n.$$

This example shows that condition (2) of theorem 4 is not sufficient without part (c), but it shows more than this. In this example parts (a) and (b) of condition (2) are satisfied and yet there are points of M which belong to no open curve of M . This raises this problem: Under what conditions does every point of a continuum M belong to some open curve of the continuum? In defining webs we assumed we were dealing with a continuous curve. With a satisfactory solution of the above problem, we might extend the definition of a web to such continua and study their properties.

IV. Concerning Webs Which are Subsets of a Continuous Curve.

10. In studying webs in general, the first question which arises is the determination of the conditions under which a continuous curve is a web. In a previous paper, „Continuous curves which are cyclically connected“¹⁸⁾, I have given five conditions under which every point of a continuous curve M belongs to an open curve of M , or in the present terminology, under which M is a web. We will insert these conditions here as

Theorem 5. In order that an unbounded continuous curve M be a web any one of the following conditions is necessary and sufficient:

¹⁸⁾ Bulletin de l'Académie Polonaise des Sciences et des Lettres, 1928, pp. 127—142.

(1) if P is any point not belonging to M and T is an inversion with center P , then $T(M) + P$ be a cyclicly connected continuous curve¹⁹⁾;

(2) if P is a cut point of M , every component of $M - P$ be unbounded;

(3) the boundary of every complementary domain of M be either a simple closed curve or an open curve [condition (3) true only in two dimensions];

(4) the M -boundary²⁰⁾ of every bounded M -domain contain more than one point;

(5) (a) M contain no end point, (b) every bounded maximal cyclic curve¹⁹⁾ C of M contain at least two limit points of $M - C$.

11. Definition. If M is a continuous curve, the set $[P]$ of all points P of M such that P lies on some open curve of M will be called the maximal web of M .

In order to justify this definition it will be necessary to show that the maximal web is actually a web. It satisfies the definition of a web providing it is a continuous curve.

Theorem 6. If M is a continuous curve, the maximal web of M is itself a continuous curve.

Proof. Let W denote the maximal web of M . If M contains no open curve, W is a vacuous set and our theorem is true vacuously. If M is the entire space, $M = W$ and our theorem is evident. If neither of these two cases holds, let P be a point of $E_n - M$ and let T be an inversion of the space with respect to P as center. The set $T(M) + P$ is a continuous curve. If X is any point of M which belongs to some open curve C of M , then $T(X)$ belongs to the simple closed curve $T(C) + P$ of $T(M) + P$. Conversely if Y is any point of $T(M)$ such that there is a simple closed curve J of $T(M) + P$ containing Y and P , then $T^{-1}(Y)$ belongs to the open curve $T^{-1}(J - P)$ of M . Hence $T(W)$ is the set of all points Q of $T(M)$ such that Q belongs to some simple closed curve of $T(M) + P$ which contains the point P , or $T(W) + P$ is the set of all simple closed curves of

¹⁹⁾ A continuous curve M is said to be cyclicly connected if every two points of M lie together on some simple closed curve of M . A cyclicly connected continuous curve C which is a subset of a continuous curve M is said to be a maximal cyclic curve of M if M contains no cyclicly connected continuous curve of which C is a proper subset. See G. T. Whyburn, *Cyclicly connected continuous curves*, Proc. Nat. Acad. Sci., vol. 13 (1927), pp. 31-38.

²⁰⁾ The M -boundary of an M -domain D of a continuum M is the set $\bar{D} - D$.

$T(M) + P$ containing the point P . By a theorem due to the author²¹), the set $T(W) + P$ is a continuous curve.

Suppose the set $T(W)$ is not connected. Then $T(W) = W_1 + W_2$, where W_1 and W_2 are non-vacuous mutually separated sets. Let Z_i be a point of W_i . Since $T(M)$ is connected, there exists an arc α of $T(M)$ with end points Z_1 and Z_2 ¹⁰). In the order from Z_1 to Z_2 let Z'_1 be the last point of W_1 on the arc α , and let Z'_2 be the first point of W_2 which follows Z'_1 . Let α' denote the subarc $Z'_1 Z'_2$ of α . There exist arcs β and γ of the continuous curves $W_1 + P$ and $W_2 + P$ with end points Z'_1 and P and Z'_2 and P respectively. Then the set $\alpha' + \beta + \gamma$ is a simple closed curve of $T(M) + P$ containing P and thus belongs to $T(W) + P$. But no point of the arc α' except its end points belongs to $T(W) + P$. Therefore $T(W)$ must be connected. Then the set $T(W)$ is connected, connected im kleinen and closed except for the point P . From this it follows that the set W is a continuous curve.

No point of the set $T(W)$, which we defined in the above proof, is a cut point of $T(W) + P$ since each point of $T(W)$ lies with P in a simple closed curve of $T(W) + P$. And we showed above that P is not a cut point of $T(W) + P$. Hence $T(W) + P$ is a cyclicly connected continuous curve²²) and it is easy to see that it is a maximal cyclic curve of $T(M) + P$. Hence we have

Theorem 7. A subset H of an unbounded continuous curve M is the maximal web of M if, and only if, either M is the entire space and $M = H$ or $T(H) + P$ is a maximal cyclic curve of $T(M) + P$ for any point P of $E_n - M$ and any inversion T with center P .

Theorem 8. If the web W is a subset of the continuous curve M , then in order that W be the maximal web of M it is necessary and sufficient that (1) no component of $M - W$ be unbounded, (2) no component of $M - W$ have more than one limit point on W .

This result is a consequence of Theorem 7 and a result due to G. T. Whyburn²³).

²¹) W. L. Ayres, On the structure of a plane continuous curve, Proc. Nat. Acad. Sci., vol. 13 (1927), pp. 749—754, theorem 2. This theorem was extended to n dimensions in my paper, Concerning continuous curves in space of n dimensions, to appear shortly in the American Journal of Mathematics.

²²) G. T. Whyburn, Cyclicly connected continuous curves, loc. cit., theorem 1. The first nine theorems of this paper by Whyburn were extended to n dimensions in my paper, Concerning continuous curves in space of n dimensions, loc. cit.

²³) Cyclicly connected continuous curves, loc. cit., theorem 2.

Theorem 9. An unbounded continuous curve M has a non-vacuous maximal web if, and only if, it does not contain a sequence of points $P_1, P_2, P_3, \dots$ such that (1) the sequence $[P_i]$ has no limit point, (2) for every i , $M - P_i$ is the sum of two non-vacuous mutually separated sets M_i^a and M_i^b , (3) for every i , M_i^b is bounded and is a subset of M_{i+1}^b .

The proof of Theorem 9 is given in my paper, „On simple closed curves and open curves“²⁴).

It may be shown that if there exists a sequence $P_1, P_2, P_3, \dots$ satisfying the conditions above, then for each point X of M there exists an integer i such that M_i^b contains the point X . Hence we have

Theorem 10. In order that an unbounded continuous curve M have a vacuous maximal web (i. e. M contains no open curve) it is necessary and sufficient that if X is any point of M there exist a cut point Y of M such that X belongs to a bounded component of $M - Y$.

Theorem 11. In order that a point P of an unbounded continuous curve M belong to the maximal web of M it is necessary and sufficient that if Q is any cut point of M , no bounded component of $M - Q$ contain P .

Proof. The condition is evidently necessary. We shall show that it is sufficient. By Theorem 10 the maximal web W of M is non-vacuous. Suppose P does not belong to W . Let L be the component of $M - W$ containing P . By Theorem 8 the set L is bounded and has just one limit point Q on W . Then Q is a cut point of M and P belongs to a bounded component L of $M - Q$. But this is contrary to the assumed condition.

Theorem 12. If M is an unbounded continuous curve, K is the set of cut points of M , and for any point X of K , M_x denotes the sum of the unbounded components of $M - X$, then

$$\Pi(M_x + X),$$

for all points X of K , is the maximal web of M .

Proof. Let $H = \Pi(M_x + X)$. If the maximal web of M is vacuous, then H is vacuous by Theorem 10. Let L be any open curve of M . If X does not belong to L , the component of $M - X$ containing L is unbounded. If X belongs to L , the two components of $M - X$ containing the two rays of $L - X$ are unbounded. Hence

²⁴) Not as yet published.

in either case L belongs to $M_x + X$. Then the maximal web of M is a subset of H . If H contains a point P that does not belong to the maximal web of M , let C be the component of M minus the maximal web of M containing P and let Q be its limit point on the maximal web of M . By Theorem 8 the point Q is a cut point of M and the component of $M - Q$ containing P is bounded. But this is impossible for P belongs to H .

12. Problem. Under what conditions does a continuous curve contain a maximal acyclic web, i. e. an acyclic web which is a proper subset of no acyclic web that belongs to the continuous curve?

(Eingegangen: 22. XI. 1928.)

Ein Theorem der Kurventheorie.

Von H. Künneth in Erlangen.

Wir nennen mit Menger eine Teilmenge M eines Raumes regulär eindimensional, wenn jeder Punkt der Menge in beliebig kleinen Umgebungen enthalten ist, deren Begrenzungen mit M endliche Durchschnitte haben. Über die regulär eindimensionalen Teilmengen kompakter Räume beweisen wir zunächst folgenden

Hilfssatz. Voraussetzung: Es liege ein kompakter Raum K vor. M sei eine regulär eindimensionale Teilmenge von K , M^* sei die Menge aller Punkte von K , in denen M eindimensional ist. p sei ein Punkt von M^* . Behauptung: Es existiert eine Umgebung $V(p)$, so daß für jede Umgebung $U(p) \subset V(p)$, deren Begrenzung B mit M einen endlichen Durchschnitt hat, die Menge B mindestens einen Punkt von $M \cdot M^*$ enthält, welcher Häufungspunkt sowohl von $M^* \cdot U(p)$, als auch von $M^* \cdot [K - \overline{U}(p)]$ ist.

Da die Menge M in p eindimensional ist, existiert eine Umgebung $V(p)$, so daß die Begrenzung von jeder Umgebung $U(p) \subset V(p)$ mit M einen nichtleeren Durchschnitt hat. Ist $U(p)$ irgend eine Umgebung $\subset V(p)$, deren Begrenzung B mit M einen endlichen Durchschnitt hat, dann behaupten wir zunächst, daß B mit $M \cdot M^*$ einen nichtleeren Durchschnitt hat. Es seien nämlich $q_1, q_2, \dots, q_m$ ($m > 0$) die Punkte von $B \cdot M$. Würde keiner dieser Punkte zu M^* gehören, so wäre in jedem von ihnen die Menge M nulldimensional. Für jeden dieser m Punkte q_i würde also eine Umgebung $U(q_i)$ mit zu M fremder Begrenzung existieren, welche zudem so klein gewählt werden könnte, daß $\overline{U}(q_i)$ den Punkt p nicht enthält. Dann wäre aber die offene Menge

$$U'(p) = U(p) - U(p) \cdot \sum_{i=1}^m \overline{U}(q_i)$$

eine Umgebung von p , die $\subset V(p)$ und deren Begrenzung zu M fremd wäre, während eine solche Menge zufolge der Bestimmung von $V(p)$ nicht existiert. Mindestens einer der m Punkte q_i von $M \cdot B$ liegt also auch in M^* . Wir haben also eine Umgebung $V(p)$ konstruiert, so daß für jede Umgebung $U(p) \subset V(p)$ die Begrenzung, falls sie mit M einen endlichen Durchschnitt hat, einen Punkt von $M \cdot M^*$ enthält.

Es sei nun $U(p)$ eine Umgebung $\subset Z(p)$, deren Begrenzung B mit M einen endlichen Durchschnitt hat. Es seien etwa $q_1, q_2, \dots, q_k$ ($k > 0$) die Punkte von $B.M.M^*$. Wir behaupten, da  mindestens einer dieser k Punkte H ufungspunkt sowohl von $M^*.U(p)$ als auch von $M^*. (K - \overline{U}[p])$ ist. W re dies nicht der Fall, so k nnte ja jedem der k Punkte eine Umgebung $U(q_i)$ zugeordnet werden, f r welche entweder die Menge $\overline{U}(q_i).U(p).M^*$ oder die Menge $\overline{U}(q_i).(K - \overline{U}[p]).M^*$ leer w re, wobei diese Umgebungen $U(q_i)$ so klein gew hlt werden k nnten, da  jede von ihnen $\subset V(p)$ ist und f r jede von ihnen die Menge $\overline{U}(q_i)$ weder p noch einen von q_i verschiedenen Punkt von $B.M$ enth lt. Dabei k nnen, da M regul r eindimensional ist, die Umgebungen $U(q_i)$  berdies so gew hlt werden, da  ihre Begrenzungen mit M einen endlichen Durchschnitt haben. Nehmen wir an, es sei etwa f r $i = 1, 2, \dots, k'$ die Menge $\overline{U}(q_i).U(p).M^*$ leer und f r $i = k' + 1, k' + 2, \dots, k$ die Menge $\overline{U}(q_i).(K - \overline{U}[p]).M^*$ leer. Dann w re

$$U'(p) = U(p) - U(p) \cdot \sum_{i=1}^{k'} \overline{U}(q_i) + \sum_{i=k'+1}^k U(q_i)$$

eine Umgebung von p , die $\subset V(p)$ w re, deren Begrenzung zu M^* fremd und als Teilmenge der Summe der Begrenzungen von $U(p)$ und den $U(q_i)$ mit M einen endlichen Durchschnitt h tte, w hrend nach dem bewiesenen ersten Teil unserer Behauptung eine solche Umgebung nicht existiert. Damit ist der Hilfssatz bewiesen.

Wir beweisen nun das

Theorem. Ist p ein Punkt des regul r eindimensionalen kompakten Raumes K , in dem die Menge N aller Punkte n -ter Ordnung von K eindimensional ist, dann ist p H ufungspunkt von Punkten von K , deren Ordnungen $\leq \frac{n}{2} + 1$ sind.

Wir wollen $p = p_1$ und $N = N_1$ setzen und bezeichnen mit N_1^* die Menge aller Punkte von K , in denen N_1 eindimensional ist. Zum Beweis des Theorems ist zu zeigen, da  jede Umgebung des Punktes p_1 einen Punkt von K von einer Ordnung $\leq \frac{n}{2} + 1$ enth lt. Es sei also $U(p)$ eine vorgelegte Umgebung. Wir haben einen Punkt von K von einer Ordnung $\leq \frac{n}{2} + 1$ in $U(p)$ aufzuweisen.

Wenden wir auf K , die Menge N_1 und den Punkt p_1 von N_1^* den Hilfssatz an, so sehen wir: Es existiert eine Umgebung $V(p_1)$, so da  f r jede Umgebung $U(p_1) \subset V(p_1)$, deren Begrenzung B endlich ist, diese Begrenzung B mindestens einen Punkt von $N_1.N_1^*$

enthält, welcher Häufungspunkt sowohl von $N_1^* \cdot U(p_1)$, als auch von $N_1^* \cdot (K - \overline{U}[p_1])$ ist. Es sei nun $U(p_1)$ eine Umgebung von p_1 , die $\subset V(p_1)$ ist und eine endliche Begrenzung B besitzt. q_1 sei ein Punkt von $N_1 \cdot N_1^* \cdot B$, der Häufungspunkt sowohl von $N_1^* \cdot U(p_1)$, als auch von $N_1^* \cdot (K - \overline{U}[p_1])$ ist. Da q_1 in N_1 liegt, existieren beliebig kleine Umgebungen von q_1 , deren Begrenzungen n Punkte enthalten. Wir wählen eine Umgebung $V(q_1)$, deren Begrenzung n Punkte enthält, so klein, daß ihr Durchmesser < 1 ist, daß $V(q_1) \subset V(p_1)$ gilt und daß $\overline{V}(q_1)$ weder p_1 noch einen von q_1 verschiedenen Punkt von B enthält. Wir betrachten nun die beiden offenen Mengen $V(q_1) \cdot U(p_1)$ und $V(q_1) \cdot (K - \overline{U}[p_1])$, welche wir bzw. V_1^i und V_1^a nennen. Ein Punkt, welcher der Begrenzung von einer dieser beiden offenen Mengen angehört, liegt entweder in der Begrenzung von $V(q_1)$ oder ist mit q_1 identisch. Zusammen enthalten also diese beiden offenen Mengen $n + 1$ Begrenzungspunkte. Der Punkt q_1 ist der einzige ihren Begrenzungen gemeinsame Punkt. Also enthält mindestens eine dieser beiden offenen Mengen V_1^a und V_1^i nicht mehr als $\frac{n}{2} + 1$ Begrenzungspunkte. Überdies enthält sowohl V_1^a als auch V_1^i einen Punkt von N_1^* , weil ja q Häufungspunkt von $N_1^*(U[p_1])$ und von $N_1^* \cdot (K - \overline{U}[p_1])$ ist.

Diejenige der beiden offenen Mengen V_1^a und V_1^i , deren Begrenzung höchstens $\frac{n}{2} + 1$ Punkte enthält (oder, falls dies für beide Mengen zutrifft, irgend eine von ihnen) bezeichnen wir mit V_1 . Sie enthält einen Punkt von N_1^* , welchen wir mit p_2 bezeichnen wollen.

Wir betrachten nun den regulär eindimensionalen kompakten Raum $\overline{V}_1$, bezeichnen mit N_2 die Menge aller Punkte n -ter Ordnung von $\overline{V}_1$ und mit N_2^* die Menge aller Punkte, in denen N_2 eindimensional ist. Wenden wir auf $\overline{V}_1$, die Menge N_2 und den Punkt p_2 von N_2^* den Hilfssatz an, so erhalten wir eine Umgebung $U(p_2)$, deren Begrenzung endlich ist und einen Punkt q_2 von $N_2 \cdot N_2^*$ enthält, welcher Häufungspunkt sowohl von $N_2^* \cdot U(p_2)$, als auch von $N_2^* \cdot (\overline{V}_1 - \overline{U}[p_2])$ ist. Durch Wiederholung des oben für q_1 durchgeführten Schlusses ergibt sich die Existenz einer offenen Menge

$V_2 \subset V_1$, deren Begrenzung höchstens $\frac{n}{2} + 1$ Punkte enthält, deren Durchmesser $< \frac{1}{2}$ ist, und welche einen Punkt von N_2^* , derselbe heiße p_3 , enthält.

In dieser Weise weiterschließend, erhalten wir eine Folge von offenen Mengen $\{V_k\}$ ($k = 1, 2, \dots$ ad inf.), die durchwegs $\subset V(p_1)$ sind, von denen jede nebst ihrer Begrenzung Teilmenge der voran-

gehenden ist, und so daß der Durchmesser von $V_k < \frac{1}{k}$ ist. Nach dem Cantorsche Durchschnitssatz existiert ein Punkt von K , welcher allen Mengen $\bar{V}_k$ gemeinsam ist; da die Durchmesser dieser Mengen gegen Null konvergieren, existiert auch nicht mehr als ein solcher Punkt und wegen $V_{k+1} \subset V_k$ für jedes k ist dieser Durchschnittspunkt auch in allen offenen Mengen V_k ($k=1, 2, \dots$ ad inf.) enthalten. Es ist also ein Punkt, der in beliebig kleinen Umgebungen enthalten ist, deren Begrenzungen höchstens $\frac{n}{2} + 1$ Punkte enthalten, d. h. ein Punkt von einer Ordnung $\leq \frac{n}{2} + 1$ von K , welcher in $V(p)$, d. h. in der vorgelegten Umgebung des vorgelegten Punktes p liegt. Die Existenz eines solchen Punktes hatten wir zum Beweise des Theorems gerade zu zeigen.

Wir bemerken noch: Aus dem Beweise des Hilfssatzes folgt, daß unter den Voraussetzungen des Hilfssatzes ein Punkt von $B.M.M^*$ Häufungspunkt nicht nur von $M^*.U(p)$ und $M^*(K - \bar{U}[p])$, sondern sogar von $M.M^*.U[p]$ und $M.M^*(K - \bar{U}[p])$ ist. Für den Beweis des Theorems ergibt sich hieraus, daß jede der beim Beweis konstruierten offenen Mengen der Folge $\{V_k\}$ einen Punkt von $N.N^*$ als inneren Punkt enthält. Infolgedessen ist der Punkt, welcher den Durchschnitt der Mengen V_k bildet, Häufungspunkt von Punkten der Ordnung n . Das Theorem kann demnach dahin verschärft werden, daß der in ihm behandelte Punkt p Häufungspunkt ist von Punkten, deren Ordnungen $\leq \frac{n}{2} + 1$ sind und die selbst Häufungspunkte von Punkten der Ordnung n sind.

Aus dem Theorem folgt als

Korollar. Eine reguläre Kurve, die in allen Punkten eine Ordnung $\leq n$ hat, enthält in jeder nichtleeren in ihr offenen Teilmenge auch einen Punkt von einer Ordnung $\leq \frac{n}{2} + 1$.

Bloß für $n=2$ kann also eine Kurve ausschließlich Punkte der Ordnung n enthalten, ein Fall, welcher bekanntlich für die Kreislinie charakteristisch ist.

(Eingegangen: 9. XI. 1928.)

Mathematik, Wissenschaft und Sprache.

Von L. E. J. Brouwer in Amsterdam.

Vortrag, gehalten in Wien am 10. III. 1928 über Einladung des Komitees zur Veranstaltung von Gastvorträgen ausländischer Gelehrter der exakten Wissenschaften.

I.

Mathematik, Wissenschaft und Sprache bilden die Hauptfunktionen der Aktivität der Menschheit, mittels deren sie die Natur beherrscht und in ihrer Mitte die Ordnung aufrecht erhält. Diese Funktionen finden ihren Ursprung in drei Wirkungsformen des Willens zum Leben des einzelnen Menschen: 1. die mathematische Betrachtung, 2. die mathematische Abstraktion und 3. die Willensauferlegung durch Laute.

1. Die mathematische Betrachtung kommt als Willensakt im Dienste des Selbsterhaltungstriebes des einzelnen Menschen in zwei Phasen zustande, die der zeitlichen Einstellung und die der kausalen Einstellung. Erstere ist nichts anderes als das intellektuelle Urphänomen der Auseinanderfallung eines Lebensmomentes in zwei qualitativ verschiedene Dinge, von denen man das eine als dem anderen weichend und trotzdem als durch den Erinnerungsakt behauptet empfindet. Dabei wird gleichzeitig das gespaltene Lebensmoment vom Ich getrennt und nach einer als Anschauungswelt zu bezeichnenden Welt für sich verlegt. Die durch die zeitliche Einstellung zustande gekommene zeitliche Zweiheit oder zweigliedrige zeitliche Erscheinungsfolge läßt sich dann ihrerseits wieder als eines der Glieder einer neuen Zweiheit auffassen, womit die zeitliche Dreiheit geschaffen ist, usw. In dieser Weise entsteht mittels Selbstentfaltung des intellektuellen Urphänomens die zeitliche Erscheinungsfolge beliebiger Vielfachheit. Nunmehr besteht die kausale Einstellung im Willensakt der „Identifizierung“ verschiedener sich über Vergangenheit und Zukunft erstreckender zeitlicher Erscheinungsfolgen. Dabei entsteht ein als kausale Folge zu bezeichnendes gemeinsames Substrat dieser identifizierten Folgen. Als besonderer Fall der kausalen Einstellung tritt auf die gedankliche Bildung von Objekten, d. h. von beharrenden (einfachen oder zusammengesetzten) Dingen der Anschauungswelt, wodurch gleichzeitig die Anschauungswelt selbst stabilisiert wird. Wie gesagt, sind die beiden Stufen der mathematischen Betrachtung keineswegs passive Einstellungen, sondern im Gegenteil

Willensakte: es kann jedermann die innere Erfahrung machen, daß man nach Willkür entweder sich ohne zeitliche Einstellung und ohne Trennung zwischen Ich und Anschauungswelt verträumen, oder die letztere Trennung aus eigener Kraft vollziehen und in der Anschauungswelt die Kondensation von Einzeldingen hervorrufen kann. Und ebenso willkürlich ist die sich nie unumgänglich aufzwingende Gleichsetzung verschiedener zeitlicher Folgen.

Die einzige Rechtfertigung der mathematischen Betrachtung ist gelegen in der „Zweckmäßigkeit“ der aus ihr hervorgehenden „mathematischen Handlung“, worunter wir folgendes verstehen. Die kausale Einstellung setzt den Menschen instand, von einer Erscheinungsfolge eine spätere, instinktiv erwünschte, aber nicht durch einen direkten Impuls herbeizuführende, als Zweck zu bezeichnende Erscheinung, indirekt durch kühle Berechnung zu erzwingen, indem man aus der Folge eine frühere, vielleicht an sich nichts begehrenswertes besitzende, als Mittel zu bezeichnende Erscheinung hervorruft, die dann die erwünschte Erscheinung als Folge nach sich zieht.

Selbstverständlich besitzt eine kausale Folge keine weitere Existenz außer als Korrelat einer mathematischen Handlungen hervorruhenden Einstellung des menschlichen Willens und kann von der Existenz eines kausalen Zusammenhanges der Welt unabhängig vom Menschen keine Rede sein. Im Gegenteil, der sogenannte kausale Zusammenhang der Welt ist eine nach außen wirkende Gedankenkraft im Dienste einer dunklen Willensfunktion des Menschen, der sich dadurch die Welt mehr oder weniger wehrlos unterwirft, in analoger Weise wie die Schlange ihre Beute wehrlos macht durch ihren hypnotisierenden Blick und der Tintenfisch durch Bespritzung mit seinem Sekret.

Die Konsequenz der kausalen Einstellung bringt weiter mit sich, daß der Mensch schon auf niedrigen Kulturstufen zur Stabilisierung seines kausalen Einflußgebietes um sich herum eine ihm untergeordnete Sphäre der Ordnung zu schaffen sucht, in welcher er erstens die ihm dienstbaren kausalen Folgen isoliert, d. h. vor störenden Nebenerscheinungen schützt, und zweitens neue kausale Folgen herbeiführt, sowohl durch die materielle Konstruktion von neuen beharrenden Objekten und Instrumenten, wie durch die mehr oder weniger organisierte Unterjochung des Willens seiner Mitmenschen unter den eigenen Willen.

2. Der volle Ausbau des Getriebes der mathematischen Handlungen wird aber erst auf höheren Kulturstufen ermöglicht, und zwar durch die mathematische Abstraktion, mittels deren man die Zweifelt ihres dinglichen Inhaltes beraubt und nur als leere Form, als gemeinsames Substrat aller Zweifelt übrig behält. Es ist dieses gemeinsame Substrat aller Zweifelt, das die Urintuition der Mathematik bildet, deren Selbstentfaltung u. a. das Unendliche als gedankliche Realität einführt und zwar in hier nicht näher zu

erörternder Weise zunächst die Gesamtheit der natürlichen Zahlen, sodann diejenige der reellen Zahlen und schließlich die ganze reine Mathematik liefert.

Die Wirkung der mathematischen Abstraktion beruht darauf, daß viele kausale Folgen erheblich leichter zu beherrschen sind, wenn man sie auf Teilsysteme derartiger reinmathematischer Systeme projiziert, d. h. ihre inhaltlosen Abstraktionen als Teilsysteme in derartige ausgedehntere reinmathematische Systeme einbettet. Hierdurch werden nämlich auch die innerhalb des ausgedehnteren Systems bestehenden Beziehungen zur Übersicht über das beschränkere System verwendbar, was für die letztere Übersicht manchmal eine durchgreifende Vereinfachung mit sich bringt. In dieser Weise kommen die wissenschaftlichen Theorien zustande, in denen neben den Elementen der kausalen Folgen das bei der Übersicht eine zentralisierende Rolle spielende erweiterte reinmathematische System als Hypothese auftritt. Speziell als exaktwissenschaftliche Theorien werden gewisse wissenschaftliche Theorien bezeichnet, die sich erstens auf ganz besonders stabile (sei es ausschließlich als Naturgesetze beobachtete, sei es als technische Tatsachen künstlich hervorgerufene) kausale Folgen beziehen, bei denen zweitens durch die Hypothesen eine große Vereinfachung erzielt wurde und bei denen drittens die kausalen Folgen speziellen Werten von zahlenmäßigen Parametern entsprechen, welche mit ihrem vollen Wertgebiet dem überlagerten mathematischen System angehören. Insbesondere bei den exaktwissenschaftlichen Theorien ereignet sich das Phänomen des heuristischen Charakters wissenschaftlicher Hypothesen, das darin besteht, daß zu ursprünglich als hypothetisch eingefügten Folgen hinterher, im überlagernden reinmathematischen System die gleiche Stelle einnehmende, wirkliche kausale Folgen der Anschauungswelt entdeckt werden.

3. Die zunächst im Dienste des Willens des einzelnen Menschen fungierende mathematische Betrachtung bzw. mathematische Handlung kann nun genau wie jede zunächst autonome aggressive oder defensive Tätigkeit als Arbeit in den Dienst eines befehlenden Willens, sei es des Einzelwillens eines anderen Menschen, sei es des Parallelwillens einer Menschengruppe oder der gesamten Menschheit, gestellt werden. Dies geschieht entweder durch als Suggestion zu bezeichnende direkte Angst- oder Schreckensanjangung, Lockung, Phantasieerregung oder animale Beherrschungskraft, oder indirekt mittels Vernunftdressur, d. h. derartige Beeinflussung der Erfahrung des dienstbar zu machenden Individuums, daß bei ihm eine, Hoffnung auf Lust oder Furcht vor Unlust als den Arbeitswillen bestimmenden Affekt auslösende, mathematische Betrachtung hervorgerufen wird.

Unter den allen Menschen vom Parallelwillen der gesamten Menschheit auferlegten mathematischen Betrachtungen ist vor allem

zu nennen die Voraussetzung der hypothetischen „objektiven Raumzeitwelt“ als gemeinsame Trägerin aller zeitlichen Erscheinungsfolgen aller Individuen; weiter die exakten und die technischen Wissenschaften, insofern sie nicht in der Form von Fabrikgeheimnissen speziellen Interessen dienen.

Als von einer beschränkteren (z. B. staatlich oder beruflich zusammengehaltenen) Menschengruppe ihren Angehörigen auferlegte mathematische Betrachtung ist in erster Linie zu erwähnen die Anerkennung und Einhaltung der Organisation der Gruppe, d. h. des Stromnetzes der Willensübertragung, mittels deren innerhalb der Gruppe die Aufzwingung der einzelnen mathematischen Betrachtungen und Handlungen als Arbeit stattfindet. Diese Organisation der einzelnen Menschengruppen hat deshalb einen viel weniger stabilen Charakter als die exakten und die technischen Wissenschaften, weil sie erstens nie alle von ihr zu berücksichtigenden äußeren materiellen Umstände beherrscht und demzufolge, um zweckmäßig zu bleiben, sich fortwährend dem Wechsel der äußeren materiellen Umstände anpassen muß, und weil zweitens ihre Effektivität nicht nur von ihrer organisatorischen Zweckmäßigkeit, sondern auch von der Treue und von der Zufriedenheit¹⁾ der ihr unterstellten Individuen abhängt, welche sich immer nur unvollkommen herbeiführen und aufrecht erhalten lassen. Denn die Treue wird vor allem an den höheren Stellen durch die Kollision der persönlichen mit den Gruppeninteressen gefährdet und die Zufriedenheit ist vor allem an den niedrigeren Stellen dadurch eine mangelhafte, daß die niedriger gestellten im allgemeinen zwar einsehen, daß die für die Angehörigen der Gruppe bestehenden gemeinsamen Wünsche und Nöte gewisse organisatorische Einrichtungen erfordern, nicht aber daß gerade die obwaltende Organisation die einzig richtige ist, und daß ihnen selbst darin die richtige Stellung zugeteilt wurde.

Um nun Treue und Zufriedenheit in den organisierten Menschengruppen, wenn auch unvollkommen, so doch leidlich zu erhalten, würden die in die Organisation aufgenommenen Mittel der Vernunftdressur bei weitem nicht ausreichen; jede Organisation ist vielmehr genötigt, überdies die Propaganda zu betreiben von moralischen Theorien, d. h. von mathematischen Betrachtungen, welche die Notwendigkeit der bestehenden Organisation außer auf egoistisch zu erfassende gemeinsame Zwecke und Nöte, überdies auf moralische, d. h. sich der egoistischen Betrachtung entziehende, Werte der Lebenshaltung zurückführen. Unmittelbar sich anbietende Beispiele sind die von der Gemeinschaft geschützten und propagierten moralischen Werte der religiösen Gebote sowie der Begriffe Vaterland, Eigentum und Familie.

Die Propaganda der moralischen Werte ist, weil sie fast keine Vernunftdressur benutzen kann, vor allem auf Suggestion,

¹⁾ Die Unzufriedenheit der einzelnen Individuen wirkt deshalb zersetzend auf die Gruppenorganisation, weil sie die Bildung von Teilgemeinschaften zeitigt, welche auf die Umformung der Organisation der Hauptgemeinschaft abzielende mathematische Betrachtungen anstellen.

insbesondere Phantasieerregung angewiesen. Übrigens beruht die Machtstellung der moralischen Werte nicht ausschließlich auf der organisierten Propaganda der entsprechenden Menschengruppe, sondern auch auf der stillen Wirkung derjenigen mathematischen Betrachtungen der einzelnen Individuen, in welche die moralischen Werte als Ablehnungen egoistischer Triebe anderer eingehen.

In den organisierten Menschengruppen kommt auf primitiven Kulturstufen und in den primitiven Beziehungen die Willensübertragung durch eine einfache Gebärde zustande und als solche ist insbesondere der Schrei effektiv. In den zur Organisation einer höheren Menschengemeinschaft gehörigen Verhältnissen dagegen sind die aufzuerlegenden Arbeiten zu verschiedenartig und zu kompliziert, um durch einfache Schreie veranlaßt werden zu können. Um die regelmäßige Veranlassung dieser Arbeiten durch bittende oder befehlende Laute zu ermöglichen, muß vielmehr die Gesamtheit der Verordnungen, Objekte und Theorien, welche bei den von den Dienstbaren verlangten mathematischen Handlungen eine Rolle spielen, selber einer mathematischen Betrachtung unterzogen werden. Den Elementen des zur aus dieser mathematischen Betrachtung erwachsenen wissenschaftlichen Theorie gehörigen reinmathematischen Systems werden sprachliche Elementarsignale zugeordnet, mit denen nach derselben wissenschaftlichen Theorie entnommenen grammatikalischen Regeln die organisierte Sprache operiert, welche die übergroße Mehrzahl der in den Kulturgemeinschaften nötigen Willensübertragungen zu bewerkstelligen erlaubt. Die Sprache ist also durchaus eine Funktion der Aktivität des sozialen Menschen. Wenn auch der einzelne Mensch in der Einsamkeit die Sprache zur Gedächtnisunterstützung braucht, so ist dies nur dem Umstande zuzuschreiben, daß er dabei die Wissenschaften und die Organisation der Gemeinschaft zu berücksichtigen hat. Und wenn auch tatlose transzendente Vorgänge von der Sprache begleitet werden, so ist dies darauf zurückzuführen, daß die gesamte menschliche Aktivität dem transzendenten Influx des freien Willens unterworfen ist.

II.

Nun gibt es aber für Willensübertragung, insbesondere für durch die Sprache vermittelte Willensübertragung, weder Exaktheit, noch Sicherheit. Und diese Sachlage bleibt ungeschmälert bestehen, wenn die Willensübertragung sich auf die Konstruktion reinmathematischer Systeme bezieht. Es gibt also auch für die reine Mathematik keine sichere Sprache, d. h. keine Sprache, welche in der Unterhaltung Mißverständnisse ausschließt und bei der Gedächtnisunterstützung vor Fehlern (d. h. vor Verwechslungen verschiedener mathematischer Entitäten) schützt. Diesem Umstande ist nicht dadurch abzuhelfen, daß man, wie es die formalistische Schule macht, die mathematische Sprache (d. h. das zur Hervorrufung reinmathematischer Konstruktionen bei anderen Menschen dienende Zeichensystem) selber einer mathematischen Betrachtung unterzieht, ihr durch Um-

arbeitung die Genauigkeit und Stabilität eines materiellen Instrumentes oder eines Phänomens der exakten Wissenschaft verleiht und sich dabei in einer Sprache zweiter Ordnung oder Übersprache über sie verständigt. Denn erstens kann beim Gebrauche der mathematischen Sprache diese Übersprache zwar mit großer Wahrscheinlichkeit (weil sie sich auf eine übersichtliche endliche Menge von beharrenden Objekten und auf die daraus abstrahierte reine Mathematik eines endlichen Systems bezieht), aber dem Wesen der Sprache entsprechend, doch nicht mit absoluter Sicherheit vor Mißverständnissen und Fehlern schützen; zweitens würde, auch wenn letzteres der Fall wäre, damit die Möglichkeit von Mißverständnissen hinsichtlich der durch eine derartige exakte mathematische Sprache angedeuteten reinmathematischen Konstruktionen keineswegs beseitigt sein.

Die Bestrebungen der formalistischen Schule, deren Ursprung nach dem obigen auf den falschen Glauben an eine magische, wenigstens an eine über ihren Charakter als Willensübertragungsmittel hinausgehende Tragweite der Sprache zurückzuführen ist, lassen sich in diesem Lichte erklären als natürliche Konsequenz eines viel älteren, primäreren, folgenschwereren und tiefer eingewurzelten Irrtums, nämlich des leichtsinnigen Vertrauens auf die klassische Logik. Dieses Vertrauen ist wie folgt entstanden: Schon im Altertum verfügte man über eine sehr vollkommene (d. h. Mißverständnisse praktisch ausschließende) Sprache der mathematischen Betrachtung von endlichen Gruppen von je als einheitlich und beharrend aufgefaßten Dingen der objektiven Raumzeitwelt. Für diese Sprache bestehen gewisse Formen des Überganges von zutreffenden (d. h. tatsächliche mathematische Betrachtungen andeutenden) Aussagen auf andere zutreffende Aussagen, welche als Gesetze der Identität, des Widerspruchs, des ausgeschlossenen Dritten und des Syllogismus bezeichnet und unter dem Namen logischer Prinzipien zusammengefaßt wurden. Wenn man diese Prinzipien rein sprachlich anwandte, d. h. mit ihrer Hilfe sprachliche Aussagen aus anderen sprachlichen Aussagen herleitete, ohne an die von diesen Aussagen angedeuteten mathematischen Betrachtungen zu denken, so erwies sich, daß sich die Prinzipien bewährten, d. h. von jeder in dieser Weise erhaltenen Aussage ließ sich hinterher konstatieren, daß sie bei jedem sprachlich erzeugenen Menschen eine tatsächliche mathematische Betrachtung auslösen konnte, welche sich für alle sprachlich erzeugten Menschen in der objektiven Raumzeitwelt praktisch als „identisch“ herausstellte.

Nun bewährten sich aber die logischen Prinzipien ebenfalls, wenn man sie ganz allgemein auf die Sprache der Wissenschaft oder auch der Begebenheiten des sonstigen praktischen Lebens kontrollierbar anwandte, wenigstens solange man dabei nur solche Ereignisse behandelte, welche von Naturgesetzen beherrscht wurden, auf deren Unerschütterlichkeit man zu vertrauen gelernt hatte. Demzufolge kam man dazu, den mittels der logischen Prinzipien hergeleiteten Aussagen auch dort zu trauen, wo sie keiner direkten Kontrolle

zugänglich waren. Insbesondere wurde auch dem Prinzip des ausgeschlossenen Dritten dieses Vertrauen entgegengebracht und dies sogar in der erweiterten Form, nach welcher ein früheres Ereignis als stattgefunden angenommen wird, nicht nur auf Grund der Absurdität, sondern auch auf Grund der praktischen Unmöglichkeit, für eine feststehende Tatsache eine andere Erklärung zu finden. Auf dieses Vertrauen stützen sich nicht nur theoretische Wissenschaften wie die Paläontologie und die Kosmogonie, sondern auch staatliche Einrichtungen wie die Strafprozeßordnung. Allerdings kam es vor, daß man in Angelegenheiten der Anschauungswelt mittels logischer Überlegungen zu falschen Ergebnissen gelangte, aber eine derartige Erfahrung hat immer nur eine geeignete Umformung der zugrunde gelegten Tatsachen bzw. Naturgesetze, nie eine Kündigung des Vertrauens auf die logischen Prinzipien zeitweilig.

Nun ist aber die Tatsache der praktischen Zuverlässigkeit der logischen Prinzipien, d. h. der Aussagenverknüpfungsgesetze der Sprache der Mathematik des Endlichen, in Angelegenheiten der Anschauungswelt nur eine Folge der allgemeineren Tatsache, daß die Menschheit die große Majorität der der Beobachtung zugänglichen Objekte und Mechanismen der Anschauungswelt in bezug auf ausgedehnte Komplexe von Tatsachen und Ereignissen erfolgreich beherrscht, indem sie das System der Zustände dieser Objekte und Mechanismen in der Raumzeitwelt als Teil eines endlichen diskreten, mit endlichvielen Verknüpfungsbeziehungen zwischen den Elementen versehenen Systems betrachtet und behandelt. M. a. W. die praktische Zuverlässigkeit der logischen Prinzipien beruht darauf, daß ein großer Teil der Anschauungswelt in bezug auf ihre endliche Organisation viel mehr Treue und Zufriedenheit zeigt als die Menschheit selbst. Daß man von altersher vor dieser nüchternen Interpretation blind war, wurde dadurch verursacht, daß man den ausschließlichen Charakter der Worte als Willensübertragungsmittel nicht erkannte und dieselben infolge eines unbesonnenen Aberglaubens als Andeutungsmittel fetischartiger „Begriffe“ betrachtete. Diese „Begriffe“ sowie die zwischen ihnen bestehenden Verknüpfungen sollten unabhängig von der kausalen Einstellung des Menschen eine Existenz besitzen, und die logischen Prinzipien sollten die Begriffe und ihre Verknüpfungen beherrschende aprioristische Gesetze darstellen. Dementsprechend herrschte die Meinung, daß Begriffsverknüpfungen, welche aus unleugbaren Axiomen (d. h. aus Begriffsverknüpfungen, welche Konstatierungen unleugbarer Tatsachen oder Naturgesetze entsprechen) mit Hilfe der logischen Prinzipien, eventuell mittels Adabsurdumführung ihres Gegenteiles, hergeleitet waren, im Falle daß sie selber wiederum kontrollierbare Aussagen über die Anschauungswelt lieferten, diese Kontrolle jederzeit siegreich bestehen könnten, im entgegengesetzten Falle aber mit gleicher Zuverlässigkeit als „ideale Wahrheiten“ zu betrachten wären. Derartige „ideale Wahrheiten“ sind denn auch Jahrhunderte lang mit zuversichtlichem Eifer von den Philosophen hergeleitet worden. Wenn die als unbehagliche Nebener-

scheinung dann und wann auftretenden Widersprüche Zweifel an der Richtigkeit dieser Entwicklungen entstehen ließen, so war dieser Zweifel nie gegen die Zuverlässigkeit der logischen Prinzipien, sondern immer nur gegen die Unleugbarkeit der Axiome, d. h. der den Entwicklungen zugrunde gelegten Begriffsverknüpfungen, gerichtet. Und manches Axiom hat man, eben auf Grund der bei den aus demselben folgenden idealen Wahrheiten auftretenden Widersprüche, verwerfen oder modifizieren müssen.

In Nachahmung der Philosophen haben schließlich auch die Mathematiker beim Studium der reinen Mathematik der unendlichen Systeme die logischen Prinzipien der Sprache der Endlichkeitsmathematik entnommen und skrupellos angewandt. In dieser Weise wurden auch für die Mathematik der unendlichen Systeme bzw. der in der Mengenlehre auftretenden, mittels des Komprehensionsaxioms geschaffenen, Mengen Aussagen „idealer Wahrheiten“ hergeleitet, welche von den Mathematikern für mehr als leere Worte gehalten wurden. Bis sich auch hier, namentlich nach Einführung der Mengenlehre, Widersprüche auftaten, und zwar solche, die sich nicht in einfacher Weise durch eine geeignete Umformung der Axiome beseitigen ließen. Diesen (in der Mathematik noch viel verblüffender als in der Philosophie wirkenden) Widersprüchen ist man zunächst mit den oben erwähnten formalistischen Bestrebungen zu Leibe gegangen. Und zwar werden hierbei unter Aufrechterhaltung des Glaubens an einen von der Willensübertragung unabhängigen Sinn der Sprache, die axiomatischen Grundverknüpfungen der mathematischen Begriffe und die zwischen den verschiedenen Verknüpfungen mathematischer Begriffe bestehenden Formen des Überganges (insbesondere sofern sie die Schaffung von Mengen und die Zulassung von Elementen zu den Mengen betreffen) einer gründlichen Analyse und Revision unterzogen, in welche selbstverständlich auch die sprachliche Wirkung der logischen Prinzipien mit hineinbezogen wird. Der Sinn der mathematischen Begriffe und Begriffsverknüpfungen wird dabei nicht näher erörtert und das Endziel der Bestrebungen (dem man allerdings noch nicht nahe gekommen ist) besteht in einer widerspruchsfreien Neugestaltung der mathematischen Sprache, die sich überdies bis auf geringe, die früheren Widersprüche umfassende, Amputationen über das ganze Lehrgebäude der bisherigen Mathematik erstrecken soll.

III.

Demgegenüber bringt der Intuitionismus die außersprachliche Existenz der reinen Mathematik zum Bewußtsein und untersucht, um auf dieser Grundlage die Richtigkeit der bisherigen Mathematik zu prüfen, zunächst, inwiefern die logischen Prinzipien, die beim Aufbau dieser Mathematik eine so große Rolle gespielt haben, auch in der Unendlichkeitsmathematik als praktisch zuverlässige Übergangsmittel zwischen reinmathematischen Konstruktionen fungieren können. Diese Untersuchung ergibt für die Prin-

zipien der Identität, des Widerspruchs und des Syllogismus ein positives, für das Prinzip des ausgeschlossenen Dritten dagegen ein negatives Resultat, d. h. es erweist sich, daß den Aussagen des letzteren Prinzips und den auf demselben beruhenden Schlußfolgerungen im allgemeinen keine mathematische Realität entspricht.

Um dies an einigen Beispielen zu erläutern, bezeichnen wir als fliehende Eigenschaft eine Eigenschaft, von der für jede bestimmte natürliche Zahl entweder die Existenz oder die Absurdität hergeleitet werden kann, während man weder eine natürliche Zahl, welche die Eigenschaft besitzt, bestimmen, noch die Absurdität der Eigenschaft für alle natürlichen Zahlen beweisen kann. Unter der Lösungszahl λ_f einer fliehenden Eigenschaft f wollen wir die (hypothetische) kleinste natürliche Zahl, welche die Eigenschaft besitzt, verstehen, unter einer Oberzahl bzw. Unterzahl von f eine natürliche Zahl, welche nicht kleiner bzw. kleiner als die Lösungszahl ist. Man sieht unmittelbar ein, daß für eine beliebige fliehende Eigenschaft jede natürliche Zahl entweder als Oberzahl oder als Unterzahl zu erkennen ist, wobei im ersteren Falle die fliehende Eigenschaft gleichzeitig ihren Charakter als solche verliert. Wir nennen die fliehende Eigenschaft f paritätsfrei, wenn man ihre Absurdität weder für die positiven noch für die negativen natürlichen Zahlen beweisen kann. Als die zur paritätsfreien fliehenden Eigenschaft f gehörige duale Pendelzahl p_f bezeichnen wir die als Limes der konvergenten Folge $a_1, a_2, \dots$ bestimmte reelle Zahl, wo a_v für eine beliebige Unterzahl v von f gleich $(-\frac{1}{2})^v$, dagegen für eine beliebige Ober-

zahl v von f gleich $(-\frac{1}{2})^{\lambda_f}$ ist. Diese duale Pendelzahl ist weder gleich Null noch von Null verschieden, im Gegensatz zum Prinzip des ausgeschlossenen Dritten. Verstehen wir unter einer nichtpositiven reellen Zahl eine reelle Zahl, die unmöglich positiv sein kann, dann ist die duale Pendelzahl weder positiv noch nichtpositiv, im Gegensatz zum Prinzip des ausgeschlossenen Dritten. Nennen wir weiter die positiven und die nichtpositiven Zahlen beide mit Null vergleichbar und die reellen Zahlen, die unmöglich mit Null vergleichbar sein können, mit Null unvergleichbar, dann ist die duale Pendelzahl weder mit Null vergleichbar noch mit Null unvergleichbar, im Gegensatz zum Prinzip des ausgeschlossenen Dritten. Und nennen wir eine reelle Zahl g rational, wenn sie entweder gleich Null ist oder zwei solche positive oder negative ganze Zahlen p und q bestimmt werden können, daß $g = \frac{p}{q}$, und irrational, wenn die Annahme der Rationalität von g ad absurdum geführt werden kann, dann ist die obige duale Pendelzahl weder rational noch irrational, im Gegensatz zum Prinzip des ausgeschlossenen Dritten.

Bezeichnen wir als die zur paritätsfreien fliehenden Eigenschaft f gehörige duale Näherungszahl n_f die als Limes der konvergenten Folge $b_1, b_2, \dots$ bestimmte reelle Zahl, wo b_ν für eine beliebige Unterzahl ν von f gleich $\left(\frac{1}{2}\right)^\nu$, dagegen für eine beliebige Oberzahl ν von f gleich $\left(\frac{1}{2}\right)^{\lambda_f}$ ist, und bringen wir in der mit einem rechtwinkligen Koordinatensystem versehenen Euklidschen Ebene eine gerade Linie l durch die Punkte $(1, p_f)$ und $(-1, n_f)$, dann sind die X -Achse und l erstens nicht parallel, während doch ihre Parallelität nicht absurd ist; zweitens fallen sie nicht zusammen, während doch ihr Zusammenfallen nicht absurd ist; drittens schneiden sie sich nicht, während doch ihr Sichschneiden nicht absurd ist.

Auch sind die X -Achse und l weder parallel, noch fallen sie zusammen, noch auch schneiden sie sich, so daß der auf dem Prinzip des ausgeschlossenen Dritten beruhende Satz, daß zwei Gerade der Euklidschen Ebene entweder parallel sind oder zusammenfallen oder aber sich schneiden, sich als hinfällig erweist.

Sollte der paritätsfreie Charakter von f verloren gehen, dann würde entweder für die Schneidung oder für die Parallelität die Gültigkeit des Prinzipes vom ausgeschlossenen Dritten zurückkehren. Aber erst wenn der Charakter von f als fliehende Eigenschaft überhaupt verloren geht, tritt das Prinzip für alle drei Eigenschaften der Parallelität, Zusammenfassung und Schneidung wieder in Kraft.

Betrachten wir das in der mit einem rechtwinkligen Koordinatensystem versehenen Euklidschen Ebene gelegene Einheitsquadrat q mit den Eckpunkten $(0, 0)$, $(0, 1)$, $(1, 0)$, $(1, 1)$. Bezeichnen wir die von q bestimmte Quadratfläche mit Q und den Punkt (p_f, p_f) mit P . Alsdann liegt P nicht auf q , während doch die Inzidenz von P mit q nicht absurd ist; weiter gehört P nicht zu Q , während doch die Zugehörigkeit von P zu Q nicht absurd ist. Schließlich gehört P weder zu q noch zum Innengebiete von q noch zum Außengebiete von q , so daß der klassische, auf dem Prinzip des ausgeschlossenen Dritten beruhende Jordansche Kurvensatz, der besagt, daß eine einfache geschlossene Kurve die Ebene in der Weise in zwei Gebiete zerlegt, daß jeder Punkt der Ebene entweder zur Kurve oder zu einem von diesen Gebieten gehört, sich als hinfällig erweist.

Betrachten wir die unendliche Reihe mit positiven Gliedern $b_1 + b_2 + b_3 + \dots$, wo die b_ν die oben angegebene Bedeutung haben. Diese Reihe konvergiert nicht, während doch ihre Konvergenz nicht absurd ist; ebenso divergiert sie nicht, während doch ihre Divergenz nicht absurd ist. Gleichzeitig erweist sich der auf dem Prinzip des ausgeschlossenen Dritten beruhende Satz, daß jede unendliche Reihe mit positiven Gliedern entweder konvergiert oder divergiert, als hinfällig. Auf diesem Satze aber, oder auf einem im wesentlichen

äquivalenten, beruht eines der wichtigsten Konvergenzkriterien aus der Theorie der unendlichen Reihen, das Kummersche Konvergenzkriterium. Und tatsächlich zeigen Gegenbeispiele, daß dieses Kriterium sich der intuitionistischen Kritik gegenüber nicht aufrecht erhalten läßt.

Betrachten wir die algebraische Gleichung $x^3 - 3x + 2b^3 = 0$, wo $b = 1 + p_f$. Die Diskriminante dieser Gleichung ist gleich $-108(1 - b^6)$, also weder gleich Null noch von Null verschieden. Auf diese algebraische Gleichung ist also der zweite Gaußsche Beweis der Wurzelexistenz nicht anwendbar. Sämtliche übrige klassische Beweise der Wurzelexistenz werden übrigens im Lichte der intuitionistischen Kritik ebenfalls hinfällig. Aber die Wurzelexistenz selber ist durch neue intuitionistische Beweise gesichert worden.

Die vorstehenden Beispiele werden verständlich machen, daß der Intuitionismus für die Mathematik weittragende Konsequenzen mit sich bringt. In der Tat müssen, wenn die intuitionistischen Einsichten sich durchsetzen, beträchtliche Teile des bisherigen mathematischen Lehrgebäudes zusammenbrechen und neue mit völlig neuem Stilgepräge errichtet werden. Und bei den Teilen, die bleiben, ist vielfach ein durchgreifender Umbau erforderlich.

Von weiteren Exkursen im Oberbau der Mathematik wollen wir aber hier Abstand nehmen und nur noch ein paar Bemerkungen grundsätzlicher Art machen. Allererst diese, daß mit dem Prinzip des ausgeschlossenen Dritten der indirekte Beweis, d. h. die Herleitung einer Eigenschaft durch *reductio ad absurdum* ihres Gegenteiles in dieser allgemeinen Form hinfällig wird. Denn die obige Pendelzahl p_f ist nicht rational, trotzdem ihre Irrationalität absurd ist, und nicht mit Null vergleichbar, trotzdem ihre Unvergleichbarkeit mit Null absurd ist. Interessant ist indessen, daß für negative Eigenschaften (d. h. Eigenschaften, die selber eine Absurdität zum Ausdruck bringen) die Methode des indirekten Beweises ungeschmälert in Kraft bleibt. Denn es gilt in der intuitionistischen Mathematik der Satz, daß Absurdität der Absurdität der Absurdität äquivalent ist mit Absurdität, so daß eine beliebige nichtverschwindende endliche Sequenz von Absurditätsprädikaten „Absurdität der Absurdität der . . . der Absurdität“, welche in der bisherigen Mathematik entweder die Richtigkeit oder die Absurdität aussagt, in der intuitionistischen Mathematik entweder mit der Absurdität oder mit der Absurdität der Absurdität äquivalent ist.

Schließlich bemerken wir noch, daß das Prinzip des ausgeschlossenen Dritten in der intuitionistischen Mathematik, obwohl nicht richtig, so doch, wenn man es ausschließlich für endliche Spezies von Eigenschaften gleichzeitig voraussetzt, widerspruchsfrei ist, was in erster Linie erklärt, daß die Irrtümer der bisherigen Mathematik sich so lange behaupten konnten, und in zweiter Linie als ermutigender Umstand für die formalistischen Bestrebungen gelten kann. Denn auf der Basis der intuitionistischen Einsichten

lassen sich außer den unabhängig vom Prinzip des ausgeschlossenen Dritten entwickelbaren richtigen Theorien, auch unter Heranziehung dieses Prinzipes mit der obigen Einschränkung, nichtkontradiktorische Theorien herleiten, mit denen sich von der bisherigen Mathematik ein viel größerer Teil als mit den richtigen Theorien umfassen läßt. Eine geeignete Mechanisierung der Sprache dieser intuitionistisch-nichtkontradiktorischen Mathematik müßte also gerade das liefern, was die formalistische Schule sich zum Ziel gesetzt hat.

Dagegen kann die gleichzeitige Aussage des Prinzips des ausgeschlossenen Dritten für beliebige Spezies von Eigenschaften sehr wohl kontradiktorisch sein. So läßt sich von der folgenden Aussage die Kontradiktorität beweisen: Alle reellen Zahlen sind entweder rational oder irrational. Im Hinblick auf diese Tatsache wird beim Aufrichten des widerspruchslosen formalistischen Sprachgebäudes doch auf jeden Fall die größte Sorgfalt und Vorsicht erforderlich bleiben.

Über dynamische Beanspruchung elastischer Systeme.

Von T. Levi-Civita in Rom.

(Vortrag, gehalten in Wien am 24. April 1928 über Einladung des Komitees zur Veranstaltung von Gastvorträgen ausländischer Gelehrter der exakten Wissenschaften.)

1. Begriffliche und mechanische Problemstellung.

Die sogenannten Resonanzerscheinungen sind längst bekannt. Es handelt sich dabei um verhältnismäßig erhebliche, manchmal sogar gefährliche, erzwungene Schwingungen, die in einem elastischen Systeme eintreten, wenn (absichtlich oder zufällig) das System unter Einwirkung periodischer Kräfte sich bewegt, deren Frequenz (genau oder annähernd) eine derjenigen ist, die man eigen nennt und die ohne periodische Anregung entstehen können. Stimmgabel und verschiedenartige Fachwerke, unter anderem Brücke und Gewölbe, bieten wichtige, häufig untersuchte Beispiele hierfür.

Jedenfalls wird, wenn elastische Körper schwingen, die Beanspruchung ihrer Kohäsion im allgemeinen größer als diejenige, welche sie unter denselben Kräften im Ruhezustand erfahren würden. Es ist offenbar wünschenswert, eine solche Vermehrung von Beanspruchung mathematisch zu erörtern und durch Formeln, die für technische Zwecke geeignet sind, abzuschätzen.

Ich werde mir erlauben, Ihnen etwas über diese ziemlich einfache, wesentlich energetische Seite der Schwingungslehre vorzutragen, die, wie wir sehen werden, sich als selbständiger Behandlung und unmittelbarer Anwendungen fähig erweist.

Es sei S irgend ein elastisches System; O seine natürliche Lage oder Konfiguration, wenn keine äußeren Kräfte einwirken. Es seien aber solche vorhanden und namentlich Konfigurationskräfte, die ausschließlich von der augenblicklichen Lage des Systems abhängen. Selbstverständlich wollen wir annehmen, daß diese äußeren Kräfte die Elastizitätsgrenzen des Systems nicht überschreiten, so daß nach den Grundlehren der Statik eine ganz bestimmte Gleichgewichtskonfiguration K , bestehen wird. In dieser Konfiguration heben sich gegenseitig die eingepprägten und die elastischen Kräfte auf. Die molekulare Kohäsion des Systems erleidet dann eine statische Beanspruchung Ω_s , die, wie man weiß, durch die entsprechende Deformationsarbeit gemessen wird: nämlich durch

die Arbeit, welche zur Überwindung der elastischen Kräfte erforderlich ist, um unser System aus seinem natürlichen Zustand O in die Gleichgewichtslage K_s (irgendwie) überzuführen.

Diese energetische Definition der Beanspruchung ist theoretisch durchaus befriedigend¹⁾ und auch von technischer Seite triftig unterstützt worden: ich beschränke mich darauf, das Werk von Castigliano²⁾ anzuführen, welches diesen Begriff systematisch verwertet.

Andererseits muß man doch erwähnen, daß in der Praxis, besonders als Kriterium für die Bestimmung der sogenannten Sicherheitsfaktoren, andere Standpunkte verbreitet sind: in diesen ist nicht die elastische Energie maßgebend, sondern, je nachdem, die größte Dehnung oder die größte Spannung oder das Verhältnis

$$\frac{\text{Schubspannung (Scherkraft),}}{\text{Normalspannung}}$$

oder auch andere derartige Bestimmungsstücke des lokalen Deformations- und Spannungszustandes.

Über diese Fragen besitzt man eine beträchtliche Literatur und ein ziemlich reiches Beobachtungsmaterial³⁾; die verschiedenen Hypothesen haben teilweise Bestätigung gefunden; man ist aber noch sehr entfernt von einer endgültigen Entscheidung.

Im großen und ganzen glaube ich, daß man folgendes festhalten kann: Innerhalb der Gültigkeit der klassischen (linearen) Elastizitätstheorie braucht man nur ein Gesamtmaß der Beanspruchung, und das wird in zutreffender Weise von der Deformationsarbeit geliefert. Dagegen bei Bruch- oder Verfließerscheinungen, wo die obigen Grenzen gewiß übertroffen sind, mag eine Gesamtabschätzung ungenügend werden und die Notwendigkeit eintreten, auch lokale Merkmale in Betracht zu ziehen.

Es sei noch einmal ausdrücklich betont, daß wir, um solche Schwierigkeiten zu vermeiden, durchaus im Gebiete der gewöhnlichen linearen Elastizitätstheorie zu verbleiben gedenken, was, überhaupt bei Strukturen und Bauwerken, sehr zu empfehlen ist.

2. Mathematische Formulierung.

Unsere Betrachtungen werden elastische Systeme aller Art umfassen; sowohl Systeme mit einer endlichen Anzahl von Freiheitsgraden, wie diskrete Massen mit elastischen Verbindungen, die als Schematisierung eines Fachwerkes gelten können, als auch kontinuierliche, wie Stäbe, Balken, Bogen, Membrane, Platten oder auch

¹⁾ Vgl. die deutlichen Überlegungen von Beltrami (Opere, T. IV, S. 180 bis 189).

²⁾ *Systèmes élastiques*, Torino, Negro, 1879.

³⁾ Vgl. etwa Love, *Mathematical theory of elasticity*, Cambridge University Press (4. Auflage), 1927, S. 121—123; oder den inhaltsreichen Enzyklopädieartikel (IV₄, 31) von Th. v. Kármán (unter Mitwirkung von L. Föppl).

massive Körper. Wir werden zwar die formelle Einkleidung des Stoffes und der Deduktionen auf Systeme mit einer endlichen Anzahl n von Freiheitsgraden beziehen. Da aber einerseits n unbestimmt bleibt und andererseits, wie wir sehen werden, die Verfahrungsweise prinzipiell und die Schlüsse durchaus unabhängig von n sind, so wird das Resultat durch unmittelbaren Grenzübergang auch für irgend welche kontinuierliche Systeme verwendbar sein.

Selbstverständlich würde es auch wohl möglich sein, die kontinuierlichen ein-, zwei- oder drei-dimensionalen Systeme jedes für sich mit denselben Begriffen, aber verschiedenem Formelapparat, je nach der Dimensionenzahl und der materiellen Struktur, zu behandeln; oder auch, wie ich ursprünglich getan hatte⁴⁾, direkt den allgemeinen Fall eines elastischen festen Körpers ins Auge zu fassen und von diesem aus (sozusagen im umgekehrten Sinne) die Grenzübergänge auszuführen, die zu Systemen mit einer oder zwei zu vernachlässigenden Dimensionen oder sogar zu diskreten Massen mit elastischen Verbindungen überzugehen gestatten.

Der erste Weg ist aber einfacher, wie das lehrreiche Muster der berühmten Sätze von Lord Rayleigh über schwingende Systeme leicht voraussehen läßt.

Es bezeichne also S irgend ein materielles System, dessen Lage durch n unabhängige Bestimmungsgrößen, d. h. Koordinaten $q_1, q_2, \dots, q_n$ festgelegt sei. Die elastische Natur des Systems sei durch die Existenz einer besondern Konfiguration O kundgegeben, die als natürlich zu bezeichnen ist, wo das System ohne äußere Kräfte in Ruhe zu beharren vermag und nach der es zurückzukehren strebt, wenn man ihm hinlänglich kleine Verschiebungen erteilt. Es ist augenscheinlich keine Beeinträchtigung der Allgemeinheit, wenn wir voraussetzen, daß die natürliche Gleichgewichtslage O den Werten

$$q_h = 0 \quad (h = 1, 2, \dots, n),$$

d. h. dem Ursprung im Konfigurationsraume entspricht.

Der wesentliche Umstand, daß das System eine innere elastische Energie besitzt, die jedem Abrücken aus O widerstrebt, läßt sich wie üblich durch die genaue mechanische Forderung deuten, daß die (nicht näher zu analysierenden) inneren elastischen Kräfte immer eine positive Arbeit leisten, wenn das System von irgend einer Konfiguration K wieder nach O übergeht. Macht man außerdem die naheliegende Annahme, daß das System an und für sich konservativ ist, so muß die eben genannte Arbeit Ω vom Wege unabhängig sein, also (da O eine feste Lage bezeichnet) muß Ω eine Funktion von K allein, d. h. der Koordinaten $q_1, q_2, \dots, q_n$ sein, die in O , d. h. für $q_h = 0$ ($h = 1, 2, \dots, n$), verschwindet: insofern Ω für alle K (mit der einzigen Ausnahme von O) positiv sein

⁴⁾ Sul massimo cimento dinamico nei sistemi elastici, Nuovo Cimento (5), II, 1901, S. 188–196.

muß, ist sie außerdem eine positiv-definite Funktion der q , die in O (d. h. für $q_1 = q_2 = \dots = q_n = 0$) das absolute Minimum Null besitzt. Ihre negativ-genommenen Ableitungen

$$-\frac{\partial \Omega}{\partial q_h} \quad (h = 1, 2, \dots, n)$$

definieren für jedes Wertsystem der q die elastischen Kraftkomponenten, welche von der Konstitution des Systems herrühren. In der Tat, wenn S von einer Lage K , die irgend welchen q -Werten entspricht, zu einer unendlich benachbarten Lage K' übergeht, die den Werten $q_h + dq_h$ entspricht, ist die Arbeit der elastischen Kräfte von K bis K' nach den vorigen Annahmen durch die Differenz darstellbar

$$\Omega(q) - \Omega(q + dq) = -d\Omega,$$

d. h. als Summe der zwei Arbeiten von K bis O und von O bis K' . Da

$$-d\Omega = -\sum_1^n \frac{\partial \Omega}{\partial q_h} dq_h,$$

so ist tatsächlich $-\frac{\partial \Omega}{\partial q_h}$ die h^{te} Lagrangesche Komponente der Kraft, wie behauptet.

Nun ist die natürliche Konfiguration O , nach unserer Annahme, eine Gleichgewichtslage und sogar (wegen des Zeichens der Arbeit von irgend einer K bis zur O) eine stabile Gleichgewichtslage für S ohne äußere Einwirkungen. Es müssen demnach die elastischen Kräfte $-\frac{\partial \Omega}{\partial q_h}$ in O verschwinden. Das hätte man übrigens schon aus der analytischen Tatsache folgern können, daß die Funktion Ω in O ein absolutes Minimum besitzt. Ihre Taylorsche Entwicklung um O beginnt somit mit quadratischen Gliedern und wir dürfen, was in elastischen Fragen ganz in der Natur der Sache liegt, uns auf einen hinreichend kleinen Bereich der q beschränken (um die natürliche Lage O), damit die Glieder dritter und höherer Ordnung ganz zu vernachlässigen sind.

Es wird demnach gestattet, einfach zu setzen

$$(1) \quad \Omega = \frac{1}{2} \sum_1^n b_{hk} q_h q_k,$$

d. h. Ω als eine positiv-definite quadratische Form der q mit konstanten Koeffizienten b_{hk} zu erblicken.

Wirken außerdem äußere Kräfte, so ist es in erster Annäherung erlaubt, ebenfalls wegen der Beschränktheit des zu erforschenden

Bereiches um O , ihre Lagrangeschen Komponenten B_h als Konstante anzunehmen. Wir werden diese äußeren Einwirkungen Belastungen nennen, da in den üblichen Problemen, welche Balken, Bogen usw. betreffen, meistens die B_h aus effektiven Belastungen stammen.

Jedenfalls (wegen der strengen oder angenäherten Konstanz der B_h) sind die Belastungen als konservative Kräfte zu betrachten, die aus der Kräftefunktion

$$(2) \quad U = \sum_{h=1}^n B_h q_h$$

herrühren.

3. Elementares Beispiel.

Es wird vielleicht sich empfehlen, einen besonders einfachen Fall, wo $n=1$ ist, vor Augen zu haben.

Man denke z. B. an einen einzigen Massenpunkt M auf vorgeschriebener Bahn, der, etwa durch eine Abrißfeder, mit einer bestimmten Stelle O der Bahnkurve verbunden ist, so daß er, aus O verrückt, eine elastische Energie

$$(1') \quad \Omega = \frac{1}{2} es^2$$

besitzt, wo e eine Elastizitätskonstante bezeichnet und s die Bogenlänge OP . Dieses s (positiv gerechnet in einem beliebig gewählten Sinne) vertritt hier die einzige Lagrangesche Koordinate. Als Belastung ist dann eine konstante tangentielle Kraftkomponente B zu erachten, so daß

$$(2') \quad U = Bs.$$

Ein noch speziellerer Fall, wo die vorgeschriebene Bahn sich einfach auf eine vertikale Linie reduziert, ist der einer schweren Masse M , die mittels einer Feder aufgehängt wird, so daß sie nur lotrechte Schwingungen ausführen kann. Sie besitzt eine elastische Energie (1'), wo s die vertikale Senkung von M bezeichnet unter die natürliche Lage O des unteren Endes der unbelasteten Feder. Die Belastung B ist hier eben das Gewicht von M .

4. Gleichgewicht unter Belastung und statische Beanspruchung.

Indem wir wieder ein allgemeines System S mit n Freiheitsgraden in Betracht ziehen, wollen wir zunächst ein paar Worte der Statik widmen.

Unter den gegebenen Belastungen B_s wird S nicht in O , sondern in einer davon verschiedenen Konfiguration K_s , wo die

Belastungen durch die elastischen Kräfte ausgeglichen werden, sich im Gleichgewichte befinden. Die dazu erforderlichen Bedingungen sind im Prinzip der virtuellen Arbeit zusammengefaßt, wonach die gesamte virtuelle Arbeit der äußeren und inneren Kräfte

$$\delta U - \delta \Omega$$

gleich Null sein muß, wenn man aus der (vorläufig unbekannten) Gleichgewichtslage K_s zu irgend einer anderen unendlich benachbarten Konfiguration K des Systems übergeht. Die Gleichung

$$(3) \quad \delta(\Omega - U) = 0$$

ergibt unmittelbar (da die unendlich kleinen Zuwächse δq_i willkürlich sind)

$$(3') \quad \frac{\partial \Omega}{\partial q_h} = \frac{\partial U}{\partial q_h} = B_h \quad (h = 1, 2, \dots, n),$$

und diese n linearen Gleichungen in den q bestimmen dieselben, d. h. K_s , eindeutig. Um dies einzusehen, hat man nur zu beachten, daß die Determinante des Systems (3') wegen (1) nichts anderes ist als die Diskriminante der quadratischen Form Ω , welche nicht verschwinden kann, da Ω positiv-definit und daher auch irreduzibel ist.

Die Gleichgewichtsbedingung (3) drückt augenscheinlich die analytische Tatsache aus, daß die Funktion $\Omega - U$ ein Extremum in K_s hat. Es liegt doch nahe zu vermuten, daß, wie für $U = 0$, die natürliche Konfiguration O ein wirkliches Minimum von Ω liefert, so dasselbe für $\Omega - U$ in der erzwungenen Gleichgewichtslage K_s geschieht. In der Tat, wegen der Linearität von U verschwindet ihr zweites Differential, so daß

$$\delta^2(\Omega - U) = \delta^2 \Omega > 0,$$

wenn nur die δq_h nicht alle gleichzeitig verschwinden. Diese Ungleichheit, zusammen mit (3), beweist gerade das Minimum von $\Omega - U$.

Wenn Gleichgewicht unter Belastung in K_s herrscht, so erfährt unser System S gemäß der Definition von Ω die statische elastische Beanspruchung Ω_s , wo offenbar Ω_s den Wert von Ω in K_s bezeichnet. Zwischen Ω_s und der entsprechenden äußeren Energie (negativ genommene Kräftefunktion) $-U_s$ der Belastungen besteht die bemerkenswerte Relation

$$(4) \quad 2\Omega_s = U_s,$$

welche sofort aus der Gleichgewichtsbedingung (3) sich entnehmen läßt.

Setzt man nämlich in (3) q_h statt δq_h (was wegen der Willkürlichkeit der δq_h erlaubt ist), so bekommt man gerade (4) nach dem Eulerschen Satze über homogene Funktionen.

Aus dem Gesagten geht der Schluß hervor: Der kleinste Wert, welchen die Differenz $\Omega - U$ überhaupt annehmen kann, ist $-\Omega_s$.

5. Dynamische Beanspruchung, die aus gegebenen Anfangsbedingungen erreichbar ist, — Definition einer oberen Grenze Ω_d — Trivialer Fall, in dem die Belastung fehlt.

Wir kommen endlich zum eigentlichen Gegenstand dieses Vortrages, nämlich zur Untersuchung der elastischen Beanspruchung im dynamischen Zustand. Vorläufig werden wir ausschließlich solche Bewegungen in Betracht ziehen, die, von beliebigem Anfangszustande ausgehend, unter Einwirkung der elastischen Kräfte und der Belastungen vor sich gehen.

Später werden wir auch die Fälle erörtern, wo anderweitige fremde Kräfte vorhanden sind, insbesondere Widerstände oder rhythmische Störungen (periodische Kräfte), vor denen man sich, wie man wohl weiß, unter Umständen sehr hüten muß. Indessen wollen wir uns mit dem einfacheren Falle eingehend beschäftigen, für welchen die engste Beschränkung erreichbar ist.

Es handelt sich dabei um Zusatzbelastungen oder augenblickliche Störungen (wie Impulse oder Stöße), welche die elastischen Schwingungen des überbelasteten oder normalen Systems hervorrufen.

Unter den durch die Anfangsbedingungen festgelegten Größen des Systems fassen wir nur ein einziges, und zwar das wichtigste Gesamtelement ins Auge, nämlich die totale Energie E , welche dem System anfangs zukommt.

Da wir jetzt rein konservative Systeme betrachten, so wird die Gesamtenergie E sich während der Bewegung erhalten. Eine solche E für unsere Systeme S besteht aus drei Anteilen: elastische Energie oder Deformationsarbeit Ω ; potentielle Energie (die von der Lage der Belastungen abhängt) $-U$; kinetische Energie oder lebendige Kraft T , die eine positiv-definite quadratische Form der Geschwindigkeiten $\dot{q} = \frac{dq}{dt}$ ist, also den Ausdruck besitzt

$$(5) \quad T = \frac{1}{2} \sum_1^n a_{hk} \dot{q}_h \dot{q}_k,$$

wo die a_{ik} im allgemeinen als Funktionen der q zu betrachten sind. Im Einklange mit der früheren Voraussetzung, daß die elastischen Vorgänge, um die es sich handelt, in einem engen Bereich um die natürliche Konfiguration O sich abspielen, würde es auch berechtigt sein, die Koeffizienten a_{ik} als Konstanten zu betrachten: doch werden wir diese letzte Annahme nicht brauchen.

Die Bewegung des Systems wird durch die Lagrangeschen Gleichungen bestimmt. Unsere Aufgabe sollte eigentlich die sein,

den maximalen Wert Ω^* zu kennzeichnen, welchen die elastische Energie Ω im Laufe der Zeit tatsächlich annimmt. Solange die Anzahl der Freiheitsgrade endlich ist, wäre das ein lösbares, wenn auch recht umständliches Problem; jedenfalls wäre das Resultat kompliziert und würde für die technisch wichtigsten, ich meine die kontinuierlichen, Systeme nicht verwendbar sein.

Glücklicherweise ist es praktisch nicht so sehr interessant, gerade Ω^* in seiner Abhängigkeit von den Anfangsbedingungen genau festzustellen, sondern es ist viel nützlicher, eine so einfache Grenze als möglich anzugeben, die Ω nicht überschreiten darf. Das gelingt tatsächlich in sehr elementarer Weise, indem man nur die Energiebedingung

$$(6) \quad T + \Omega - U = E$$

berücksichtigt und das Maximum Ω_d von Ω , welches mit (6) verträglich ist, aufsucht. Selbstverständlich hat man

$$\Omega^* \leq \Omega_d,$$

weil Ω^* nicht nur durch (6) (welche eine notwendige Folge der Bewegungsgleichungen ist), sondern durch alle übrigen Bewegungsgleichungen bedingt ist.

Ehe wir auf die Untersuchung von Ω_d eingehen, wird es sich empfehlen zu bemerken, daß, wenn es sich um eigentliche Bewegung (nicht etwa Ruhe) unseres Systems handelt, die Ungleichung bestehen wird

$$(7) \quad E > -\Omega_s.$$

In der Tat ist die Konstante E (in jedem Augenblicke und insbesondere zu Anfang) gleich $T + \Omega - U$; T , als lebendige Kraft, darf nie negativ sein, und das zweite Glied $\Omega - U$ besitzt, wie wir im vorigen Paragraph gesehen haben, ein absolutes Minimum in der Gleichgewichtslage K_s , nämlich, gemäß (4),

$$\Omega_s - U_s = -\Omega_s.$$

Da wir diesen Gleichgewichtszustand gerade ausschließen, so wird zu Anfang entweder $T > 0$, oder die Lage des Systems nicht mit K_s übereinstimmen, und daher (7) befriedigt sein.

Was die eigentliche Bestimmung von Ω_d angeht, sind wir auf eine ganz gewöhnliche Aufgabe der Differentialrechnung zurückgeführt.

Man hat $2n$ Veränderliche: die n Koordinaten q_h und die n Geschwindigkeiten $\dot{q}_h$, die durch die einzige Bedingungsgleichung (6) verbunden sind. Es handelt sich darum, das (relative) Maximum einer Funktion $\Omega(q_1, q_2, \dots, q_n)$ der q allein zu entnehmen.

Im trivialen Falle, wo die Belastung fehlt ($U=0$), ist die Aufgabe unmittelbar erledigt. In der Tat, die Gleichung (6) reduziert sich dann auf

$$T + \Omega = E,$$

wo T und Ω beide ≥ 0 sind. Der größte Wert Ω_a von Ω ist demnach erreicht für $T=0$ und ist somit

$$(8) \quad \Omega_a = E.$$

6. Bestimmung von Ω_a im allgemeinen Fall ($U \neq 0$) — Sicherheitskoeffizient.

Zuvörderst bemerken wir in bezug auf das gesuchte Maximum von Ω , daß die Werte der Argumente $q, \dot{q}$ durch keine Ungleichung beschränkt sind. Eventuelle Grenzmaxima sind somit ausgeschlossen, dagegen müssen die gewöhnlichen Stationaritätsbedingungen notwendigerweise erfüllt sein. Nach der bekannten Lagrangeschen Regel ist dementsprechend

$$(9) \quad \delta(\Omega + \lambda E) = 0$$

zu setzen, wo λ einen vorläufig unbestimmten konstanten Multiplikator bezeichnet, unter E ist das erste Glied von (6) zu verstehen, und unter δ ein Differentiationssymbol, welches willkürlichen, unendlich kleinen Zuwächsen der Argumente q und $\dot{q}$ entspricht.

Wenn man berücksichtigt, daß die Argumente $\dot{q}$ nur in T vorkommen, so entnimmt man zunächst aus (9)

$$(10) \quad \lambda \frac{\partial T}{\partial q_h} = 0 \quad (h = 1, 2, \dots, n),$$

und folglich, indem man mit $\frac{1}{2} \dot{q}_h$ multipliziert und die n Gleichungen mit einander addiert,

$$(10') \quad \lambda T = 0.$$

Es darf nicht vorkommen, daß $\lambda = 0$, denn, gemäß (9), würde dann $\delta\Omega = 0$, was einem Extremum von Ω entsprechen würde; und ein solches Extremum kann, wegen der Definition (1) von Ω , nur ein Minimum (namentlich Null) sein; und wir fragen nach einem Maximum.

Da $\lambda \neq 0$ sein muß, so verlangt (10') $T = 0$. T ist eine positiv-definite quadratische Form der $\dot{q}$, die nur dann verschwindet, wenn

alle $\dot{q}_h = 0$ ($h = 1, 2, \dots, n$). Dies bringt das Verschwinden aller $\frac{\partial T}{\partial \dot{q}_h}$ mit sich, womit tatsächlich den Bedingungen (10) Genüge geleistet wird.

Es ist also, in (9), $T = 0$ zu setzen, so daß wegen (6) die Bedingung die Form annimmt:

$$(1 + \lambda) \delta \Omega - \lambda \delta U = 0.$$

Der Fall $U = 0$ wurde früher durch (8) erschöpft; U ist somit als $\neq 0$ anzusehen. Infolgedessen ist es gestattet, vorige Gleichung durch $1 + \lambda$ zu dividieren, weil die Annahme $1 + \lambda = 0$, $\delta U = 0$ nach sich ziehen würde und daher, wegen (2), das Verschwinden aller Belastungen B_h , also von U selbst.

Die Stationaritätsbedingung (9) läßt sich folglich schreiben:

$$(9') \quad \delta \Omega - \mu \delta U = 0,$$

wo $\mu = \frac{\lambda}{1 + \lambda}$ den ursprünglichen Multiplikator λ ersetzt, und, wie jener, eine vorläufig unbestimmte (von Null verschiedene) Konstante bezeichnet, die nachträglich durch die Bedingung (6) (für die stationären Werte) zu bestimmen ist.

Wir sind jetzt imstande, unsere kleine Rechnung begrifflich zu verwerten. Es genügt nämlich (9') mit (3) zu vergleichen, um einzusehen, daß (9') gerade eine Gleichgewichtslage unseres Systems bestimmt, und zwar diejenige, die der Belastung μU entspricht. Nun erhält man, falls man in den Gleichungen (3') (die linear in den q sind) μU statt U einführt,

$$(11) \quad \frac{\partial \Omega}{\partial q_h} = \mu \frac{\partial U}{\partial q_h} = \mu B_h,$$

so daß die zu bestimmenden Größen $q_h^{(d)}$ (Koordinaten dieser fiktiven Gleichgewichtslage) mit $\mu q_h^{(s)}$ übereinstimmen: selbstverständlich bezeichnen die $q_h^{(s)}$ die Koordinaten der Gleichgewichtslage K_s , d. h. die Lösungen von (3').

Daraus folgt einfach

$$(12) \quad \Omega_d = \mu^2 \Omega_s,$$

$$(13) \quad U_d = \mu U_s,$$

wo d auf den stationären, s auf den statischen Zustand sich bezieht. Gemäß der Gleichung

$$(4) \quad 2\Omega_s = U_s$$

ergibt sich aus (12) und (13)

$$\Omega_d - U_d = (\mu^2 - 2\mu) \Omega_s.$$

Die Konstante μ ist so zu bestimmen, daß die stationären Werte $[T_d = 0, (12) \text{ und } (13)]$ die Energiebedingung (6) befriedigen. Dies fordert, gemäß der letztgefundenen Gleichung,

$$(14) \quad (\mu^2 - 2\mu) \Omega_s = E.$$

Beachtet man noch einmal, daß wir den besonderen Fall $U = 0$ (als schon durch (8) erledigt) jetzt ausschließen, so ist man berechtigt (da K_s gewiß nicht mit 0 zusammenfällt), Ω_s als positiv zu betrachten.

Die Gleichung für μ läßt sich somit schreiben

$$(14') \quad (\mu - 1)^2 = 1 + \frac{E}{\Omega_s}.$$

Die linke Seite ist, gemäß der Ungleichung (7), jedenfalls > 0 , und die Gleichung (14) hat somit zwei reelle Wurzeln.

Die eine

$$(14'') \quad \mu = 1 + \sqrt{1 + \frac{E}{\Omega_s}},$$

wo dem Wurzelzeichen sein arithmetischer Wert beizulegen ist, ist > 1 . Die andere dagegen < 1 und hat überhaupt einen kleineren absoluten Wert.

Um die Maximumeigenschaft des stationären Wertes $\Omega_d = \mu^2 \Omega_s$ zu prüfen, ist nur die Bestimmung (14'') von μ in Erwägung zu ziehen, da die andere Wurzel augenscheinlich einen kleineren Wert von Ω ergeben würde.

Dies festgestellt, fassen wir wieder die Energiebedingung (6) ins Auge und bemerken einerseits, daß, im stationären Zustand, $T_d = 0$,

$$\Omega_d - U_d = E;$$

andererseits daß, in irgend einem Zeitpunkte t , während der Bewegung,

$$T + \Omega - U = E \text{ ist.}$$

Durch Subtraktion bekommt man

$$(15) \quad \Omega_d - \Omega + (U - U_d) = T.$$

Wendet man die Taylorsche Entwicklung auf die quadratische Form $\Omega(q_1, q_2, \dots, q_n)$ in der Umgebung des Wertsystems $q_h^{(d')}$ an, so kann man schreiben

$$\Omega - \Omega_d = \sum_1^n \left(\frac{\partial \Omega}{\partial q_h} \right)_d (q_h - q_h^{(d)}) + \Omega',$$

wo Ω' nichts anderes als Ω selbst, berechnet für die Argumente $q_h - q_h^{(d)}$, bezeichnet, und daher ≥ 0 ist.

Nun darf man wegen (11) unter dem Summenzeichen statt $\left(\frac{\partial \Omega}{\partial q_h} \right)_d$ einfach μB_h einführen, so daß die Summe selbst sich auf $\mu(U - U_d)$ reduziert.

Hieraus haben wir nach Division durch μ

$$\frac{1}{\mu} (\Omega - \Omega_d) = U - U_d + \frac{\Omega'}{\mu}.$$

Addiert man zu (15), so ergibt sich

$$\left(1 - \frac{1}{\mu} \right) (\Omega_d - \Omega) = T + \frac{\Omega'}{\mu}.$$

Das zweite Glied kann nie negativ werden; wegen $\mu > 1$, ist der Faktor $1 - \frac{1}{\mu}$ zweifellos positiv. Unter diesen Umständen zeigt die obere Gleichung, daß tatsächlich

$$\Omega_d - \Omega \geq 0$$

bestehen muß; das Gleichheitszeichen gilt nur, wenn gleichzeitig $T = 0$, $\Omega' = 0$, d. h. wenn das System ev. die Lage $q_h^{(d)}$ mit Geschwindigkeit Null einnimmt.

So haben wir mit aller Strenge bewiesen, daß dem Ausdruck

$$\Omega_d = \mu^2 \Omega_s$$

die Eigenschaft einer oberen Grenze der dynamischen Beanspruchung zukommt, und es ist wohl gerechtfertigt, den Faktor $\mu^2 (> 1)$, mit welchem die statische Beanspruchung Ω_s zu multiplizieren ist, um die obere Grenze der dynamischen zu erhalten, als Sicherheitskoeffizienten zu bezeichnen.

7. Besprechung des Resultats — Beispiele.

Wenn man in

$$(13) \quad \Omega_d = \mu^2 \Omega_s$$

für μ seinen Wert

$$(14'') \quad \mu = 1 + \sqrt{1 + \frac{E}{\Omega_s}}$$

einsetzt, so nimmt Ω_d die explizite, aber wenig übersichtliche Form an

$$(13') \quad \Omega_d = 2 \Omega_s + E + 2 \sqrt{\Omega_s(\Omega_s + E)},$$

welche, gegenüber (12) zusammen mit (14''), den einzigen Vorzug hat, nicht der Beschränkung $U \neq 0$ unterworfen zu sein, sondern allgemein zu gelten auch im Grenzfalle, wo die Belastungen fehlen und demnach $\Omega_s = 0$ ist. Um dies einzusehen, hat man bloß zu beachten, daß (13') für $\Omega_s = 0$ sich auf

$$\Omega_d = E$$

reduziert, in Übereinstimmung mit (7).

Es wird nunmehr angemessen sein, das Ergebnis durch einige Beispiele zu erläutern.

Betrachten wir zunächst den öfters angegebenen Fall, wo gegebene Belastungen (von Kräftefunktion U) zu wirken anfangen, in dem das System sich in O und Ruhezustand befindet. Wäre das System wieder im Gleichgewicht unter dieser Belastung, so würde es die statische Beanspruchung Ω_s erfahren. Aber bevor das Gleichgewicht sich einstellt, bewegt sich das System. Seine Anfangsbedingungen sind folgende: $T = 0$, $\Omega = 0$, $U = 0$, also $E = 0$. Wir entnehmen somit aus (14'') $\mu = 2$ und bekommen als obere Grenze der Beanspruchung während der Einsetzungsschwingungen

$$\Omega_d = 4\Omega_s.$$

Die dynamische Beanspruchung kann somit viermal so groß als die entsprechende statische werden.

Es ist übrigens auch leicht festzustellen, indem man sich auf den einfachen Fall eines einzigen Freiheitsgrades bezieht, daß die obige Grenze tatsächlich erreicht werden kann.

Fassen wir dafür die Schwingungen einer Feder ins Auge, nachdem man an dem unteren Ende ein Gewicht B angehängt hat.

Wenn man die Masse der Feder in Vergleich zu $m = \frac{B}{g}$ vernachlässigt, so bekommt man aus

$$(1') \quad \Omega = \frac{1}{2} es^2,$$

$$(2') \quad U = Bs,$$

$$(5') \quad T = \frac{1}{2} m \dot{s}^2$$

die (Lagrangesche) Bewegungsgleichung

$$\ddot{s} = -\omega^2 s + g,$$

wo, der Kürze halber, ω^2 statt $\frac{e}{m}$ geschrieben ist.

Diese Gleichung (die man natürlich auch in der elementarsten Weise gewinnen kann) besitzt die partikuläre Lösung

$$s = \frac{g}{\omega^2},$$

welche dem Gleichgewichte in der Lage, wo die Belastung durch die elastische Kraft kompensiert wird, entspricht.

Die elastischen Schwingungen, d. h. das allgemeine Integral, erhält man, indem man zu der partikulären Lösung g/ω^2 das Integral der homogenen Gleichung

$$\ddot{s} = -\omega^2 s$$

addiert. Es handelt sich dabei um harmonische Schwingungen um 0 mit der Frequenz $\omega/2\pi$, deren Ausdruck

$$s = r \cos (\omega t + \theta_0)$$

ist, wo $r (>0)$ und θ_0 Integrationskonstante bezeichnen. Sie sind so zu bestimmen, daß im ersten Augenblicke ($t=0$) Ruhe herrscht in $s=0$. Aus

$$s = r \cos (\omega t + \theta_0) + g/\omega^2$$

erhält man somit

$$s = \frac{g}{\omega^2} (1 - \cos \omega t)$$

und die größte Abweichung aus der Gleichgewichtslage beträgt

$$2 g/\omega^2$$

(für ωt gleich einem ungeraden Vielfachen von π). Da die Beanspruchung $\Omega = \frac{1}{2} e s^2$ mit s^2 proportional ist, so nimmt sie tatsächlich (für $s = 2 g/\omega^2$) einen Wert an, der viermal so groß ist als die statische Bestimmung, welche $s = g/\omega^2$ entspricht.

Gehen wir jetzt zu einer ein wenig allgemeineren Anwendung über, indem wir voraussetzen, daß gerade vor der Einwirkung der Belastungen B_h das System sich im Gleichgewichte befand unter Einwirkung anderer Belastungen, die in konstantem Verhältnisse α

zu den B_h sind. Wenn $\alpha > 1$, hat man offenbar mit einer plötzlichen Reduktion der Belastung (im Verhältnis $\frac{1}{\alpha}$) zu tun, wie sie z. B. infolge eines partiellen Losgehens derselben vorkommen kann.

Wenn dagegen α einen echten Bruch bedeutet, so hat man mit einer Vermehrung der Belastung (immer im Verhältnis $\frac{1}{\alpha}$) zu tun, was dem Eintreten von Zusatzbelastungen entspricht. Endlich bedeutet negatives, daß außerdem der Sinn der Kräfte umgekehrt wird.

Wie dem auch sei, die Gleichungen (12) und (13), in denen μ durch α zu ersetzen ist, zeigen, daß neben $T=0$ als Anfangswerte

$$\Omega = \alpha^2 \Omega_s, \quad U = \alpha U_s$$

heranzuziehen sind.

Man hat folglich, indem man sich erinnert, daß $2\Omega_s = U_s$ ist,

$$E = T + \Omega - U = (\alpha^2 - 2\alpha)\Omega_s,$$

$$1 + \frac{E}{\Omega_s} = (1 - \alpha)^2.$$

Der absolute Wert von $\sqrt{1 + \frac{E}{\Omega_s}}$ ist demnach $\alpha - 1$ für $\alpha \geq 1$, dagegen $1 - \alpha$ für $\alpha \leq 1$, worin auch der Fall α negativ mitinbegriffen ist. Mithin ergibt sich aus (14'')

$$\mu = \begin{cases} \alpha, & \text{falls } \alpha \geq 1 \\ 2 - \alpha, & \text{falls } \alpha \leq 1 \end{cases}$$

und aus (12)

$$\Omega_d = \begin{cases} \alpha^2 \Omega_s \\ (1 - \alpha)^2 \Omega_s \end{cases}$$

je nachdem $\alpha \geq 1$ oder ≤ 1 ist.

Es ist zu bemerken, daß die elastische Beanspruchung Ω , welche nach (1) eine quadratische Funktion der q ist, auch als eine quadratische Funktion (und zwar die reziproke Form) der durch (3') definierten zugehörigen Belastungen B angesehen werden kann. Man erkennt somit, daß $\alpha^2 \Omega_s$ nichts anderes als die ursprüngliche elastische Beanspruchung bedeutet. Wir haben daher die interessante (und nur halb anschauliche) Folge, daß im dynamischen Zustande, welcher von einer Reduktion der Belastungen hervorgerufen wird, keine Vermehrung der ursprünglichen statischen Beanspruchung zu befürchten ist.

Treten dagegen Zusatzbelastungen ein in derselben Richtung wie die ursprünglichen, was $\alpha < 1$ und > 0 entspricht, so ist das

Maß der zu befürchtenden Zunahme durch den Sicherheitskoeffizienten

$$(15) \quad \mu^2 = (2 - \alpha)^2$$

gegeben, woraus erhellt, daß beim Eintreten von Zusatzbelastungen die zu befürchtende dynamische Beanspruchung während der Einstellungsperiode höchstens viermal so groß ist als diejenige, die sich schließlich im Gleichgewichte einstellt. Der Wert 4 von μ wird übrigens nur für $\alpha = 0$ genommen, wie das vorige Beispiel (anfängliche Ruhe im natürlichen Zustande) wiedergibt. Besteht dagegen, auch vor der Einwirkung von Zusatzkräften, eine nicht verschwindende Belastung, so ist $\alpha > 0$, und das bewirkt $\mu^2 < 4$ (jedoch > 1).

Endlich, wenn die Veränderung der Belastungen mit Umkehrung des Sinnes begleitet ist ($\alpha < 0$), dann ist μ^2 immer > 4 , und kann jeden noch so hohen Wert erreichen, wenn $|\alpha|$ genügend groß ist. Besonders zu erwähnen ist der Fall $\alpha = -1$, wo eine einfache Umkehrung aller Belastungen ohne Änderung der Intensität stattfindet. Man hat dann $\mu^2 = 9$.

Wie im ersten speziellen Beispiele, würde es, indem man sich wieder auf Systeme mit einem einzigen Freiheitsgrade bezieht, leicht sein zu bestätigen, daß die angegebenen maximalen Beanspruchungen tatsächlich erreicht werden können.

8. Die Beanspruchung als Funktion von E — Allgemeine Wirkung von Widerstandskräften.

Aus dem Ausdruck (14'') von μ springt in die Augen, daß μ immer mit E wächst (im zulässigen Variationsintervall dieser Größe, welches, nach § 5, $-\Omega_s \vdash \infty$ ist). Hieraus folgt, daß auch die dynamische Beanspruchung $\Omega_d = \mu^2 \Omega_s$ eine immer wachsende Funktion von E ist.

Dies festgestellt, wenden wir uns zu der den reellen Verhältnissen entsprechenden Voraussetzung, daß neben elastischen Kräften und Belastungen noch passive Widerstände, wie Reibung, Viskosität, usw. vorhanden sind. Ohne sie zu analysieren, werden wir uns begnügen, die naheliegende, aber wichtige Tatsache hervorzuheben, daß, genau wie in der Statik, das Vorkommen von dissipativen Kräften (d. h. von Kräften, deren Arbeit immer negativ ist) zu Gunsten der Sicherheit wirkt, in dem Sinne, daß die größtmögliche dynamische Beanspruchung gewiß nicht höher ausfällt als diejenige, die, caeteris paribus, im konservativen Fall (d. h. wenn keine Störung vorhanden wäre) zu erwarten ist.

Der Beweis ist sehr einfach. Man hat nur statt (6) die allgemeinere Energie-Bilanz heranzuziehen, wo auch die Arbeit $-\Psi$ der dissipativen Kräfte (von Anfang bis zur Zeit t) berücksichtigt ist.

Die Gleichung, welche die obere Grenze von Ω bedingt, ist somit

$$(16) \quad T + \Omega - U = E - \Psi,$$

wo Ψ , seiner Definition nach, eine nie negative Größe ist, die für $t = 0$ verschwindet.

Im (wenn auch unbekannten) Zustande der maximalen elastischen Beanspruchung wird (16) erfüllt sein mit einem (gleichfalls unbekannten, aber ≥ 0) Werte Ψ_a von Ψ . Nennen wir der Kürze halber E' die $E - \Psi_a$, so hat man das Maximum von Ω zu bestimmen, welches mit

$$T + \Omega - U = E'$$

verträglich ist. Das ist aber genau unsere frühere Aufgabe (§§ 5—6) für dasselbe System unter ausschließlich konservativen Kräften, wenn den Anfangsbedingungen eine (totale) Energie E' zukommt, die zwar unbekannt, aber nicht größer als E ist.

Das entsprechende Maximum Ω'_a wird dann gewiß $\leq \Omega_a$, w. z. b. w.

Im allgemeinen darf man nicht mehr behaupten; und es ist auch leicht, sich zu überzeugen, daß unter passenden Anfangsbedingungen, trotz der Einwirkung von Widerständen, die obere Grenze Ω_a tatsächlich erreicht wird.

Doch hat neulich Herr Krall⁵⁾ die interessante Frage näher erörtert und (durch Heranziehung von Normalkoordinaten) eine bemerkenswerte Klasse elastischer Systeme mit innerer Reibung angegeben, für deren Schwingungen (aus der Ruhe) sich ein Sicherheitskoeffizient $\mu' < \mu$ angeben läßt. Genauer findet man

$$\mu' - 1 = (\mu - 1) e^{-\gamma/\nu},$$

wo γ eine von der Natur des Materials (Elastizität und innere Reibung) abhängende Konstante bezeichnet und ν die Grundschnitzungszahl der freien Schwingungen ist.

9. Lemma.

Es seien, wie üblich, mit Q_h bezeichnet die Lagrangeschen Komponenten irgend welcher unserem Systeme eingepprägten Kräfte.

Bekanntlich hängen, wenn ihre lokale Verteilung, d. h. die Vektoren F_i , welche die in einzelnen Punkten angreifenden Kräfte darstellen, gegeben sind, die Q_h von den F_i folgendermaßen ab:

$$Q_h = \sum_i F_i \cdot \frac{\partial P_i}{\partial q_h},$$

⁵⁾ Rend. della R. Accademia dei Lincei, Vol. VII, 1928, S. 223—228, 556—561.

wo die Summe auf alle materiellen Punkte P_i des Systems zu erstrecken ist, und $\cdot$ ein skalares Produkt bezeichnet.

Wie dem auch sei, die elementare Arbeit, welche die Q_h leisten, wenn die Koordinaten q_h die elementaren Zuwächse $dq_h = \dot{q}_h dt$ erfahren, hat den Ausdruck

$$dt \sum_1^n Q_h \dot{q}_h,$$

so daß die elementare Wirkung w und die Gesamtarbeit A während eines endlichen Zeitintervalles $0 \rightarrow \tau$ bzw. sind

$$(17) \quad w = \sum_1^n Q_h \dot{q}_h,$$

$$(18) \quad A = \int_0^\tau w dt.$$

Es wird sich empfehlen, neben den (kovarianten) Kraftkomponenten Q_h die entsprechenden kontravarianten Q^h einzuführen, die durch die Gleichungen

$$(19) \quad Q_h = \sum_1^n a_{hk} Q^k$$

definiert sind: wo selbstverständlich die a_{hk} , wie in (5), Koeffizienten der lebendigen Kraft T bezeichnen. Der Ausdruck (17) von w läßt sich dann auch schreiben

$$(17') \quad w = \sum_1^n a_{hk} \dot{q}_h Q^k.$$

In den folgenden Überlegungen wird auch die Größe

$$(20) \quad \gamma = \frac{1}{2} \sum_1^n Q_h Q^h = \frac{1}{2} \sum_1^n a_{hk} Q^h Q^k$$

erscheinen, welche als statischer Zwang zu bezeichnen ist: (20) ist nämlich eine (invariante, das heißt von der Wahl der Lagrangeschen Parameter q unabhängige) positive Funktion der gegebenen Kräfte, welche mit dem klassischen Gaußschen Zwang eng zusammenhängt⁶⁾. In der Tat, für freie Punktsysteme ohne Beschleunigung deckt sich (20) genau mit der Gaußschen Definition

$$\gamma = \frac{1}{2} \sum_i \frac{1}{m_i} F_i^2,$$

⁶⁾ Vgl. etwa Levi-Civita und Amaldi, *Lezioni di meccanica razionale*, Vol. II, p. 487 (Bologna, Zanichelli, 1926).

wo m_i die Masse des Punktes P_i bedeutet; jedenfalls besitzt (20) die Eigenschaft, daß die Gleichgewichtsbedingungen $Q_h = 0$, die mit

$$Q^h = 0 \quad (h = 1, 2, \dots, n)$$

gleichbedeutend sind, mit den Minimumsbedingungen der Funktion γ übereinstimmen.

Wir wollen jetzt eine wichtige Ungleichung ableiten, die den absoluten Wert von A beschränkt. Es handelt sich im wesentlichen um eine Schwarzsche Ungleichung, die sich folgendermaßen beweisen läßt.

Wegen des positiv-definiten Charakters der Form T , hat man notwendigerweise

$$\frac{1}{2} \int_0^\tau dt \sum_{h,k} a_{hk} (\dot{\xi} \dot{q}_h + \eta \dot{q}_k) (\dot{\xi} \dot{q}_h + \eta \dot{q}_k) \geq 0,$$

für irgend welche reelle Werte der Argumente ξ und η . Betrachtet man diese Argumente als von t unabhängig, so kann die vorstehende Ungleichung, vermöge (5), (17'), (18) und (20), auch so geschrieben werden:

$$\xi^2 \int_0^\tau T dt + 2\xi\eta A + \eta^2 \int_0^\tau \gamma dt \geq 0,$$

woraus die gewünschte Beschränkung für A^2 und somit auch für $|A|$ folgt:

$$(21) \quad A^2 \leq \int_0^\tau T dt \cdot \int_0^\tau \gamma dt.$$

Eine ähnliche Ungleichung gilt für irgend einen Wert τ^* zwischen 0 und τ , insbesondere für denjenigen Augenblick, in dem die (stetige) Funktion $|A|$ ihr Maximum A^* erreicht. Es besteht somit neben (21)

$$A^{*2} \leq \int_0^{\tau^*} T dt \cdot \int_0^{\tau^*} \gamma dt.$$

Da T und γ wesentlich positiv sind, werden die Integrale des zweiten Gliedes gewiß nicht vermindert, wenn man beide bis τ (statt bis τ^*) erstreckt. In dieser Weise erhält man

$$(21') \quad A^{*2} \leq \int_0^\tau T dt \cdot \int_0^\tau \gamma dt,$$

welche die frühere Ungleichung (21) in sich umfaßt.

Führt man die Mittelwerte

$$\overline{\gamma} = \frac{1}{\tau} \int_0^\tau \gamma \, dt, \quad \overline{T} = \frac{1}{\tau} \int_0^\tau T \, dt$$

des Zwanges und der lebendigen Kraft, im betrachteten Intervalle, ein, so läßt sich (21') auch schreiben:

$$(21'') \quad A^{*2} \leq \tau^2 \overline{\gamma} \overline{T}.$$

10. Obere Schranke der dynamischen Beanspruchung unter beliebigen Kräften.

Wenn auf unser System S , außer den elastischen Kräften und den (konstanten) Belastungen B_h , noch die allgemeinen Q_h einwirken, so nimmt offenbar, unter Berücksichtigung von (18), die Energiegleichung (6) die Gestalt [(16), wo A an Stelle von $-\Psi$ zu schreiben ist] an:

$$(22) \quad T + \Omega - U = E + A.$$

Bedenkt man einerseits (§ 4), daß $-\Omega_s$ das Minimum von $\Omega - U$ liefert, so daß

$$U - \Omega \leq \Omega_s,$$

und andererseits daß

$$A \leq |A| \leq A^*,$$

so folgt aus (22) die Ungleichung

$$T \leq E + \Omega_s + A^*,$$

welche in jedem Augenblicke t des Zeitintervalles $0 \leq t \leq \tau$ gültig ist. Durch Integration in bezug auf t zwischen 0 und τ und Division durch τ erhält man ferner (da E , Ω_s und A^* alle konstant sind)

$$(23) \quad \overline{T} \leq E + \Omega_s + A^*.$$

Benutzt man, in (21''), diese Ungleichung für $\overline{T}$, so ergibt sich

$$(24) \quad A^{*2} \leq \tau^2 \overline{\gamma} \{ E + \Omega_s + A^* \},$$

was mit

$$\left\{ A^* - \frac{1}{2} \tau^2 \overline{\gamma} \right\}^2 \leq \tau^2 \overline{\gamma} \left\{ E + \Omega_s + \frac{1}{4} \tau^2 \overline{\gamma} \right\}$$

gleichbedeutend ist, oder endlich mit

$$(24') \quad A^* \leq \tau W,$$

wo, der Kürze halber,

$$(25) \quad W = \frac{1}{2} \tau \bar{\gamma} + \sqrt{\bar{\gamma} \left\{ E + \Omega_s + \frac{1}{4} \tau^2 \bar{\gamma} \right\}}$$

gesetzt ist.

Diese Konstante W hängt, wie man sieht, von 4 Argumenten ab: die vielfach gebrauchte Anfangsenergie E und statische Beanspruchung Ω_s (die überhaupt nur in der Kombination $E + \Omega_s$ vorkommen); die Dauer τ des Intervalles, um das es sich handelt; und der (mittlere) Zwang $\bar{\gamma}$ während desselben Intervalles.

Bedenkt man, daß (mit den üblichen Maxwell'schen Bezeichnungen) die Dimensionen einer Energie und daher von E , Ω_s und A^* , $l t^{-2} m$ sind, so entnimmt man aus (24), daß $\tau^2 \bar{\gamma}$ dieselbe Dimension besitzt, so daß

$$[\bar{\gamma}] = l^2 t^{-4} m,$$

wie übrigens auch direkt aus seiner Definition als Zwang leicht folgt.

Die numerische Berechnung von $\bar{\gamma}$ ist unmittelbar ausführbar im besonderen, aber praktisch sehr wichtigen Fall, in dem die Kräfte Q_h bekannte, namentlich periodische oder sogar sinoidale, Funktionen der Zeit sind.

Jedenfalls erhellt aus (24'), daß die obere Grenze A^* von A die Dauer τ als Faktor enthält, während der andere Faktor W auch für $\tau \rightarrow 0$ endlich bleibt.

Nach dieser Vorbereitung sind wir imstande, an die Gleichung (22) dieselben Überlegungen anzuknüpfen, die wir früher (§ 8) auf die Gleichung (16) angewandt haben.

Es genügt, die veränderliche Größe A durch die obere Grenze A^* ihrer absoluten Werte zu ersetzen, und so kommt man wieder (für die Bestimmung einer oberen Grenze von Ω im Intervall $0 \leq t \leq \tau$) auf das in den §§ 5—6 gelöste Problem.

Es gilt somit, von $t=0$ bis $t=\tau$,

$$\Omega_d = \mu^2 \Omega_s,$$

wo, im Ausdrucke (14'') von μ , die Energie E durch $E + \tau W$ zu ersetzen ist: W ist durch (25) gegeben.

Gemäß (20) und (22) verlangt die Berechnung von $\bar{\gamma}$ und mithin von W nur Quadraturen, wenn die Q_h explizit als Funktionen der Zeit gegeben sind.

Wenn das System zu Anfang (vor der Einwirkung der Q_h) sich im Gleichgewichte unter den Belastungen B_h befand, hat man speziell für $t=0$, $T=0$, $\Omega=\Omega_s$, $U=U_s=2\Omega_s$, also gemäß (6)

$$E + \Omega_s = 0.$$

Der Ausdruck (25) von W vereinfacht sich dann erheblich und reduziert sich auf

$$(25') \quad W = \tau \bar{\gamma}.$$

Dementsprechend wird der Wert (14'') von μ (indem man statt E $E + \tau W = -\Omega_s + \tau^2 \bar{\gamma}$ einführt)

$$(26) \quad \mu = 1 + \tau \sqrt{\frac{\bar{\gamma}}{\Omega_s}}.$$

Wir werden im nächsten Absatze ein paar Beispiele anführen. Indessen wollen wir noch eine Bemerkung qualitativer Natur hinzufügen.

Es ist allgemein bekannt, daß, wenn auf ein elastisches System, welches, wie unsere S , freie Schwingungen auszuführen fähig ist, noch periodische Kräfte wirken, deren Periode nahezu gleich einer der freien ist, dann gefährliche Resonanzeffekte zu befürchten sind, d. h. die Schwingungsweiten und somit die Beanspruchung Ω sehr beträchtlich werden, sobald der erzwungene periodische Zustand erreicht (oder nahezu erreicht) wird.

Das soeben ausgesprochene Resultat zeigt doch, daß die Erhöhung der Beanspruchung nur allmählich ist: für τ hinreichend klein, kann sie nicht bedeutend höher werden als diejenige Ω_d , welche ohne störende Einflüsse im allerschlimmsten Fall eintritt.

11. Einige Formeln, die die Durchbiegung eines an den Enden gestützten Balkens betreffen.

Es sei S ein Balken der Länge l , dessen linker Stützpunkt O wie gewöhnlich als Koordinatenanfangspunkt genommen wird, während man Ox als horizontal und mit der Achse des Balkens in seiner natürlichen Lage zusammenfallend annimmt; Oy sei vertikal nach unten gerichtet. Wenn der Balken sich durchbiegt, bleibt seine Achse in der vertikalen Oxy -Ebene und kann durch eine Gleichung der Form

$$y = \sum_1^n q_h \sin \frac{\pi h x}{l}$$

wiedergegeben werden, in der die q_h (deren Zahl unendlich ist) die Lagrangeschen Koordinaten des vorliegenden Falles darstellen.

Eine in einem Punkte des Balkens mit der Abszisse x konzentrierte Last hat gemäß der in § 9 gegebenen Formel folgende Lagrangesche Komponenten

$$Q_h = Y \frac{\partial y}{\partial q_h} = Y \sin \frac{\pi h x}{l} \quad (h = 1, 2, \dots).$$

Handelt es sich dagegen um eine Last mit dem Gewicht p' pro Längeneinheit, die gleichförmig auf eine Strecke $a \text{ --- } b$ des Intervalls $0 \text{ --- } l$ verteilt ist, so hat man analog

$$(27) \quad Q_h = p' \int_a^b \sin \frac{\pi h x}{l} dx = \frac{p' l}{\pi h} \left(\cos \frac{\pi h a}{l} - \cos \frac{\pi h b}{l} \right).$$

Aus der Definition von T folgt, wenn wir der Einfachheit halber von dem Teil absehen, der aus der Rotation der Querschnitte entspringt, sofort

$$(28) \quad T = \frac{p}{2g} \int_0^l \dot{y}^2 dx = \frac{pl}{4g} \sum_1^\infty \dot{q}_h^2,$$

wo p/g die lineare Dichte des Balkens bedeutet; aus der Definition von Ω ⁷⁾ hat man ferner

$$(29) \quad \Omega = \frac{1}{2} B \int_0^l \left(\frac{d^2 y}{dx^2} \right)^2 dx = \frac{B}{4} \frac{\pi^4}{l^3} \sum_1^\infty h^4 q_h^2,$$

wo $B = EJ$ die sogenannte Biegesteifigkeit, E den Youngschen Elastizitätsmodul und J das Trägheitsmoment des Balkenquerschnittes bezüglich der neutralen Achse bedeutet.

Die Gleichungen der freien Bewegung

$$\frac{d}{dt} \frac{\partial T}{\partial q_h} - \frac{\partial (T - \Omega)}{\partial q_h} = 0 \quad (h = 1, 2, \dots)$$

nehmen die Form an

$$\ddot{q}_h + h^2 \omega^2 q_h = 0,$$

wo

$$(30) \quad \omega^2 = \frac{\pi^4 g B}{p l^4}$$

ist. Es handelt sich also um Schwingungen mit den Perioden $2\pi/h\omega$. Die Fundamentalperiode T entspricht dem Wert $h = 1$ und gemäß (30) ist ihr Ausdruck gegeben durch

⁷⁾ Vgl. etwa Rayleigh, The theory of sound, Vol. I, S. 257 (2. Auflage, London, Macmillan, 1894).

$$(30') \quad T^2 = \frac{4 p l^4}{\pi^2 g B}.$$

Wir beschränken uns auf die Betrachtung an den Enden gestützter Balken in den beiden typischen Fällen statischer Beanspruchung, in denen vorliegt:

1. Ein gleichförmig verteiltes Gewicht der Größe p pro Längeneinheit.

2. Eine Last mit dem Gewicht P , die in einem Punkt der Abszisse ξ konzentriert ist.

Für Ω_s ergeben sich bzw. die Ausdrücke⁸⁾

$$(31) \quad \Omega_s = \frac{1}{240} \frac{p^2 l^5}{B},$$

$$(32) \quad \Omega_s = \frac{1}{96} \frac{P^2}{l B} \xi^2 (l - \xi)^2.$$

Was speziell (31) angeht, so ist es von Nutzen, zwei transformierte Ausdrücke anzugeben: den einen, indem man an Stelle von B die Fundamentalperiode T durch (30') einführt, wodurch man erhält

$$(31') \quad \Omega_s = \frac{1}{240} \frac{1}{4} \pi^2 g p l T^2$$

und den andern, indem man die vertikale elastische Verschiebung des Balkenmittelpunktes oder den Biegunspfeil einführt, dessen Wert (Biegunspfeil)

$$f = \frac{1}{24} \cdot \frac{5}{16} \cdot \frac{p l^4}{B}$$

ist; damit erhält man

$$(31'') \quad \Omega_s = \frac{16}{50} P f,$$

worin P die Gesamtlast $p l$ bezeichnet.

Im Falle der Gleichung (32), d. h. einer einzigen, im Punkte ξ konzentrierten Last, hat man hingegen

$$\Omega_s = \frac{1}{2} P \eta,$$

worin η die elastische Verschiebung des Punktes ist, in dem die Last angebracht ist.

⁸⁾ Siehe z. B. A. E. H. Love, The mathematical theory of elasticity, S. 371—372 (4. Auflage, Cambridge University Press, 1927).

12. Der Eisenbahnzug auf der Brücke.

Wenn ein Eisenbahnzug, dessen mittleres Gewicht pro Längeneinheit p' sei, über eine Brücke fährt, so ist die Brücke in einem beliebigen Augenblicke einer zusätzlichen Inanspruchnahme unterworfen, die einer Belastung p' pro Längeneinheit in einer variablen Strecke $a \text{ — } b$ des Segments $0 \text{ — } l$ entspricht. Nehmen wir der Einfachheit halber an, daß der Zug mit der Geschwindigkeit v im Sinne der x -Achse läuft und die Zeitählung in dem Augenblicke beginnt, in dem die Durchlaufung der Brücke beginnt, so wird in einer ersten Phase am Anfang der Brücke ein Segment $0 \text{ — } b$ belastet, wo $b = v \cdot t$ ist, und zwar solange, bis $v \cdot t$ die kleinere der beiden Längen: $l = \text{Länge der Brücke}$ oder $l_1 = \text{Länge des Zuges}$ erreicht; ist dann $l_1 < l$, so folgt die Belastung einer variablen Strecke $a \text{ — } b$; hingegen wird $0 \text{ — } l$ belastet, wenn $l_1 > l$ ist usw. Jedenfalls entsprechen die Belastungsbedingungen immer dem Schema (27), wo a, b stetige Funktionen von t sind, die zwischen 0 und l liegen.

Aus dem Ausdruck (28) von T folgt, auf Grund von (19),

$$Q_h = \frac{2g}{pl} Q_h,$$

woraus auf Grund der Definition (20) des statischen Zwanges γ , wenn man noch n statt h schreibt,

$$(33) \quad \gamma = \frac{2g}{pl} \sum_1^\infty Q_n^2$$

sich ergibt. Man beachte nun, daß in jeder der Reihen

$$s_a = \sum_1^\infty \frac{1}{n^2} \cos^2 \frac{\pi n a}{l}, \quad s_b = \sum_1^\infty \frac{1}{n^2} \cos^2 \frac{\pi n b}{l},$$

$$s' = \sum_1^\infty \frac{1}{n^2} \cos \frac{\pi n a}{l} \cos \frac{\pi n b}{l}$$

das allgemeine Glied einen absoluten Wert $\leq \frac{1}{n^2}$ hat. Da nun $\sum_1^\infty \frac{1}{n^2}$

konvergent mit der Summe $\frac{\pi^2}{6}$ ist, so konvergieren die hingschriebenen Reihen gleichmäßig bezüglich a und b und haben jede eine Summe $\leq \frac{\pi^2}{6}$.

Die (notwendig positive) Größe

$$s_a + s_b - 2s'$$

ist also $\leq \frac{4}{6} \pi^2 = \frac{2}{3} \pi^2$, und man kann daher setzen

$$s_a + s_b - 2s' = \frac{2}{3} k(t) \pi^2,$$

worin $k(t)$ einen echten Bruch, und zwar eine stetige Funktion von t bezeichnet.

Daher folgt aus (27) und (33)

$$\gamma = \frac{4}{3} g \frac{p'^2}{p} l k(t),$$

und nimmt man den Mittelwert im Intervalle $0 \text{ --- } \tau$, so kommt

$$(34) \quad \bar{\gamma} = \frac{4}{3} g \frac{p'^2}{p} l \bar{k},$$

worin $\bar{k}$ ebenso wie k nicht die Einheit überschreitet.

Wir haben nunmehr alles beisammen, um die obere Grenze der dynamischen Beanspruchung, die dem Passieren eines Zuges, d. h. der vereinten Wirkung der vergrößerten Belastung und der in der Brücke hervorgerufenen Schwingungen entspricht, zu berechnen. Es genügt zu dem Zwecke, sich den entsprechenden Sicherheitskoeffizienten μ zu verschaffen, dessen Ausdruck

$$(26) \quad \mu = 1 + \tau \sqrt{\bar{\gamma} / \Omega_s}$$

wir dem Paragraphen 10 entnehmen können, da die Brücke vor der Überfahrt des Zuges sich im Gleichgewicht unter alleiniger Wirkung ihres eigenen Gewichtes befindet. Für Ω_s nehmen wir den Ausdruck (31') und man hat daher auf Grund von (34)

$$\bar{\gamma} / \Omega_s = \bar{k} \frac{240 \cdot 4^2}{3} \frac{1}{\pi^2} \left(\frac{p'}{p} \right)^2 \frac{1}{T^2}.$$

Fassen wir die numerischen Faktoren in einen einzigen Koeffizienten zusammen und setzen für $\bar{k}$ seinen größtmöglichen Wert 1 ein (womit man, was wohl selbstverständlich ist, die Sicherheit nur erhöht), so folgt

$$\tau \sqrt{\bar{\gamma} / \Omega_s} = 11,38 \frac{p'}{p} \frac{\tau}{T}.$$

Daher erhält man schließlich für den Sicherheitskoeffizienten

$$(26') \quad \mu = 1 + 11,38 \frac{p'}{p} \frac{\tau}{T},$$

worin p das Gewicht der Brücke, p' das der Last pro Längeneinheit, τ die Gesamtdauer der Überfahrt des Zuges über die Brücke und T die Fundamentalperiode der Schwingungen der Brücke bedeuten.

Eine solche obere Grenze, die unter den von uns gemachten Schematisierungen absolut sicher ist, ist in Wahrheit in den gewöhnlichen Fällen nicht so vorteilhaft, in denen die üblichen Abschätzungen einen weniger hohen Wert von μ plausibel machen⁹⁾. Während aber diese üblichen Kriterien ein gefährliches Anwachsen der Beanspruchung vermuten lassen für den Fall, daß, bei zunehmender Geschwindigkeit der Züge, τ dem Werte der freien Eigenschwingungsperiode der Brücke T nahekommt, zeigt unsere Formel (26'), daß diese Vermutung unbegründet ist; im Gegenteil, die Zunahme der Geschwindigkeit hat eine Vergrößerung der Sicherheit zur Folge.

Schließlich wollen wir auch den Einfluß kleiner Unebenheiten der Gleise untersuchen. Wenn δ die maximale Niveaudifferenz bedeutet, die von irgend einer Spalte in den Gleisen herrührt, so erhält eine punktförmige Last P' durch Fallen oder Steigen die lebendige Kraft $P'\delta$. Es ist natürlich nicht gesagt, daß eine solche lebendige Kraft ganz auf die darunter befindliche Brücke übertragen wird, im Gegenteil, im allgemeinen wird nur ein geringer Bruchteil sich wirklich in Energie der Vertikalschwingungen des Trägers umsetzen. Jedenfalls handelt man im Sinne einer Erhöhung der Sicherheit, wenn man annimmt, daß die gesamte Energie $P'\delta$ auf die Brücke übertragen wird. Und da es sich nur um Abschätzungen handelt, ist es auch erlaubt, denselben Ausdruck beizubehalten, auch wenn P' nicht eine konzentrierte Last bedeutet, sondern die Summe der Lasten, die den Sprung δ machen, wie dies bei einem Eisenbahnzuge des Gesamtgewichtes P' der Fall ist.

Den Einfluß der kleinen Unebenheit δ kann man, wenn man ihn isoliert betrachtet, d. h. so als ob die lebendige Kraft $P'\delta$ der Brücke beigebracht würde, während sie sich unter ihrem eigenen Gewicht im Gleichgewicht befindet, der Gleichung (14'') entnehmen, indem man darin durch

$$E = T + \Omega - U$$

den Wert der Energie einführt, der den Anfangsbedingungen

$$T = P'\delta, \quad \Omega = \Omega_s, \quad U = U_s$$

entspricht, d. h., wegen (4),

$$E = -\Omega_s + P'\delta.$$

⁹⁾ Vortreffliche Angaben hierfür findet man in dem interessanten Artikel von S. Timoshenko, *Vibrations of bridges*, American Society of mechanical Engineers (Annual meeting, New York, 1927).

Es folgt daraus:

$$\mu = 1 + \sqrt{P'\delta/\Omega_s}$$

und schließlich, indem man für Ω_s seinen Wert (31'') einsetzt und dabei noch $\frac{16}{50}$ in $\frac{1}{3}$ abrundet,

$$\mu = 1 + \sqrt{\frac{P'}{P} \frac{\delta}{3f}},$$

wo P das Gesamtgewicht der Brücke, P' das der beweglichen Last, δ den Sprung und f den Biegunspfeil bedeuten, der ausschließlich dem eigenen Gewichte der Brücke entspricht. Diese Formel ersetzt mit Vorteil eine analoge Abschätzung, die in der Technik üblich ist und durch unsichere Betrachtungen gewonnen wird¹⁰⁾.

¹⁰⁾ Vgl. etwa C. Guidi, *Lezioni sulla scienza delle costruzioni*. Parte II (11. Auflage, Torino, Bona, 1925), Kap. V, Art. 168—171.

Literaturberichte.

E. Fettweis, Das Rechnen der Naturvölker. Bei Teubner, Leipzig und Berlin 1927. 8°, IX + 89 Seiten. RM 5,—.

Dem Vorwort entnimmt man als Ziel der Arbeit eine Darstellung und Analyse des Rechnens der Naturvölker, um dadurch einerseits der Entscheidung gewisser strittiger Fragen der psychologisch begründeten Rechenmethodik zu dienen, andererseits einer Korrektur gewisser Urteile über die Rechenkunst und in Anschluß daran über die geistige Veranlagung der Naturvölker. Man findet mit großer Sorgfalt die Nachrichten über Naturvölker aller Weltteile außer Europa zusammengetragen, systematisch geordnet und mit der methodischen und psychologischen Literatur in Beziehung gesetzt. Das Ergebnis ist in acht Thesen zusammengefaßt, von denen die letzte für den rechtzeitigen Fortschritt von der Sinnlichkeit zur Abstraktion eintritt, im Gegensatz zu entgegengesetzten Tendenzen in der neuesten Rechenmethodik.

Von den mitgeteilten Tatsachen sei eine erwähnt (§ 7). Die Savaras in Südindien rechnen nach einem Zwölfersystem und erklären, das komme daher, daß eines Tages, als Angehörige ihres Stammes auf dem Felde über zwölf hinauszählen wollten, ein Tiger sie alle aufgefressen hätte.

Darin liegt vielleicht der Hinweis auf die Schwierigkeit und Verwirrung, welche eintritt, wenn man ohne äußere Hilfsmittel über gewisse Grenzen im Rechnen hinausgeht. Ein ausführlicheres Literaturverzeichnis erhöht das Interesse an der Studie.

Wirtinger.

O. Neugebauer, Die Grundlagen der ägyptischen Bruchrechnung. Bei Springer, Berlin 1926. RM 7,50.

Der Verfasser bezeichnet als wichtigstes prinzipielles Ergebnis seiner Arbeit die Einsicht in die ausschließlich additive Grundlage der ägyptischen Mathematik und stützt dieses Ergebnis auf eingehende Betrachtung der ägyptischen Bruchrechnung, wie sie im Papyrus Rhind überliefert ist.

Wirtinger.

Die Kegelschnitte des Apollonios, übersetzt von A. Czwalina. Bei R. Oldenbourg, München und Berlin 1926. RM 10,—.

Die Übersetzung beschränkt sich auf die ersten vier Bücher, welche allein griechisch überliefert sind. Hoffentlich trägt sie dazu bei, den berühmten Geometer nicht bloß zu zitieren, sondern auch zu lesen. Hervorgehoben sei, daß das vierte Buch — obgleich vielleicht am wenigsten durchgearbeitet, die Frage nach der Anzahl der Schnittpunkte zweier Kegelschnitte behandelt und damit bewußt über die sonst übliche Beschränkung auf Probleme zweiten Grades hinausgeht.

Die Anmerkungen beschränken sich auf das Nötigste und belasten das Buch nicht. Ebenso ist auf textkritische und andere Weitläufigkeiten verzichtet, so daß der klassische Meister unmittelbar auf den Leser wirken kann. Solche Ausgaben sind sehr zu begrüßen.

Wirtinger.

Briefwechsel zwischen Carl Friedrich Gauss und Christian Ludwig Gerling. Herausgegeben im Auftrage der Gesellschaft zur Beförderung der gesamten Naturwissenschaft zu Marburg von Dr. Clemens Schaefer. Mit einem Bildnis von Chr. Gerling und einem Faksimilebrief von Schweikart. XX + 820 Seiten. Gr. 8°. Bei O. Elsner, Berlin 1927. Geh. RM 35,—.

Die 400 Jahrfeier der Universität Marburg gab den Anlaß zu der sehr sorgfältigen Herausgabe dieses Briefwechsels, der sich in 388 Briefen von dem Jahre 1810 bis zum Jahre 1854 erstreckt. Ein Register und recht sorgfältige Anmerkungen machen die Sammlung noch wertvoller.

Wenn auch heute dem Bilde von Gauss kaum sehr wesentliche Züge hinzugefügt werden können, so ist doch auch für ihn eine Reihe von Einzelheiten von Interesse.

Als Beispiele seien angeführt Gauss' eigene Mitteilung über die Auffindung der Konstruktion des Siebzehneckes, seine Bemerkungen über den Begriff des Flächeninhaltes sich selbst schneidender ebener Figuren, die Bemerkungen über Kongruenz und Symmetrie und rechts und links gewundene Schrauben, weiterhin die Bemerkung (S. 676, 8. IV. 1844) über eine Geometrie von mehr als drei Dimensionen, wofür die menschlichen Wesen keine Anschauung haben, die aber in abstracto nicht widersprechend sei und füglich höheren Wesen zukommen könne.

Eine Seite 783 f. beschriebene, Gauss eigentümliche Anordnung, um den Einfluß der Erdrotation auf die Schwingungsebene des Pendels in kurzer Zeit sichtbar zu machen, scheint ebenfalls wenig oder gar nicht bekannt zu sein.

In den Seite 673 f. beschriebenen Versuchen Bunsens und Gerlings vom 21. III. 1844 kann man die ersten Anfänge des Scheinwerfers sehen.

Endlich berichtet Gerling ausführlich über die Versuche über das Tischrücken und Gauss antwortet mit der Erzählung eigener Versuche. Beide sind durchaus skeptisch und Gauss erwähnt mit bei ihm seltenen Humor einen Heidelberger Professor, der, von Figur ein wahrer Falstaff, dem davonlaufenden Tisch gar nicht nachkommen konnte.

Was etwa sonst an interessanten, beobachtungstechnischen, die Ausgleichung betreffenden, sowie geodätischen, astronomischen und erdmagnetischen Fragen in dem Briefwechsel berührt wird, werden Kundigere an der Hand des Registers leicht herausfinden.

Es ist gewiß eine wichtige Quelle zur Kultur- und Gelehrtengegeschichte der Zeit in sorgfältigster Bearbeitung vorgelegt und insbesondere Gerling damit ein würdiges Denkmal gesetzt.

Wirtinger.

D. Hilbert und W. Ackermann, Grundzüge der theoretischen Logik. (Die Grundlehren der mathematischen Wissenschaften in Einzeldarstellungen, Band XXVII.) Bei Springer, Berlin 1928. VIII und 120 Seiten. RM 7,60, geb. RM 8,80.

Dieses Buch ist aus Vorlesungen Hilberts entstanden, die von Bernays ausgearbeitet und nun von Ackermann definitiv redigiert wurden. Das erste Kapitel bringt den Aussagenkalkül, zunächst inhaltlicher, dann in axiomatischer Darstellung, wobei auch die Verbesserungen zur Geltung kommen, die Bernays an dem in den Principia Mathematica von Whitehead und Russell benützten Axiomensysteme angebracht hat. Wesentlich über den bisher üblichen Rahmen der Theorie geht hinaus der Nachweis der Widerspruchslosigkeit, Unabhängigkeit und Vollständigkeit des Axiomensystemes. Unverständlich blieb dem Referenten, wozu in § 2 neben der Beziehung der Gleichwertigkeit ($x \sim y$), die besagt, daß x und y gleichen Wahrheitswert haben, noch eine Beziehung der Äquivalenz ($x \dot{\sim} y$) eingeführt wird, die besagen soll, daß x und y „gleichbedeutend“ sind, — ein dem Aussagekalkül fremder Begriff, der dementsprechend auch alsbald wieder aus der Darstellung verschwindet. Das zweite Kapitel (Der Prädikaten- und Klassenkalkül) handelt ganz kurz von Aussagefunktionen $\varphi(x)$ einer Veränderlichen und bringt insbesondere die Einordnung der Aristotelischen Schlußformen in den Kalkül. Das dritte Kapitel (Der engere Funktionenkalkül) handelt von Aussagefunktionen mit beliebig vielen Veränderlichen, doch werden nur All- und Seinszeichen zugelassen, die sich auf Individuen beziehen, nicht solche, die sich auf Funktionen beziehen. Ausgegangen wird auch hier von einem Axiomensystem von Bernays, das gegenüber dem der Principia Mathematica eine Verbesserung bedeutet. Die Widerspruchslosigkeit wird bewiesen und mit Ackermann gezeigt, daß Vollständigkeit im strengen Sinne hier nicht besteht; die Unabhängigkeit ist noch nicht untersucht. Besonders sei noch hingewiesen auf die sehr einfache Lösung des Entscheidungsproblems für den Fall, daß nur Aussagefunktionen mit einer Veränderlichen auftreten, der dem vereinigten Aussagen- und Klassenkalkül, also der klassischen Aristotelischen Logik, entspricht. Im vierten Kapitel (Der erweiterte Funktionenkalkül) werden auch All- und Seins-

zeichen zugelassen, die sich auf Funktions- und Aussagenvariable beziehen. Es werden die Paradoxien vorgeführt, die sich hier bei naivem Gebrauch der Symbolik einstellen und die Lösung dieser Paradoxien durch Russells Typentheorie wird kurz besprochen. Es kommen auch die Schwierigkeiten zur Sprache, die sich auf dem Boden der Typentheorie einem logischen Aufbau der Mathematik entgegenstellen, und die Beseitigung dieser Schwierigkeiten durch das Reduzibilitätsaxiom. — Hilberts Aufbau der Mathematik soll in einem späteren Bande dargelegt werden. — Am Schlusse dieser ausgezeichneten, übersichtlichen und leicht lesbaren Darstellung der theoretischen Logik findet man ein ausführliches Literaturverzeichnis.

Hans Hahn.

A. Fraenkel, Einleitung in die Mengenlehre. (Grundlehren der mathematischen Wissenschaften, Band 9.) Dritte Auflage. Bei Springer, Berlin 1928. RM 22,60, geb. RM 24,—.

Daß nach verhältnismäßig kurzer Zeit eine dritte Auflage dieses Buches notwendig geworden ist, ist das deutlichste Zeichen für seine Beliebtheit. Wenn heute Cantors Grundgedanken auch in weiteren nichtmathematischen Kreisen bekannt und diskutiert zu werden beginnen, so ist dies sicher zum großen Teil ein Erfolg von Fraenkels ausgezeichnete Einleitung in die Mengenlehre.

Die außerordentlich klar und faßlich geschriebenen drei ersten Kapiteln (Über die Grundlagen, die Kardinalzahlen, Ordnungstypen und Ordinalzahlen) sind im wesentlichen aus der zweiten Auflage übernommen. Sie wurden erweitert um Übungsaufgaben am Ende einiger Abschnitte und, vielleicht etwas zu reichlich, in den Zitaten. Eine starke Umarbeitung hat die zweite Hälfte des Buches betreffend die Grundlagenfragen erfahren.

Das Kapitel IV („Erschütterung der Grundlagen und ihre Folgen“) beginnt mit einer Erörterung der Paradoxien, an welche sich eine Darstellung des Intuitionismus schließt. Fraenkel bezeichnet als die beiden Hauptpunkte der Gedanken Brouwers erstens seine Ansicht über die Unabhängigkeit der Mathematik (des mathematischen Konstruierens) von der Logik (der Sprache dieser Konstruktionen) und die damit zusammenhängende Ablehnung des Satzes vom ausgeschlossenen Dritten und zweitens die Aufstellung einer konstruktiven Mengendefinition. Während der erste Punkt von Fraenkel erschöpfend behandelt wird, läßt auch die neue Auflage eine eingehende Darstellung des zweiten Punktes vermissen. Es unterbleibt selbst eine Angabe von Brouwers Begriffsbestimmung, die „schwer verständlich und heute noch nicht als allseits geklärt zu betrachten“ sei (S. 237). Der Referent glaubt, in einer (übrigens erst nach Drucklegung von Fraenkels Buch erschienenen) Note (Jahresbericht d. D. M. V. 37, S. 213) die Mengendefinition, welche Brouwer an die Spitze seiner Entwicklung der Mengenlehre unabhängig vom Satz vom ausgeschlossenen Dritten stellt, dem Verständnis leicht zugänglich gemacht zu haben: Sie ist, bei Verwendung des Satzes vom ausgeschlossenen Dritten, äquivalent mit der Definition der analytischen Mengen, welche man z. B. in der zweiten Auflage von Hausdorffs Mengenlehre (unter dem Namen Suslinsche Mengen) eingehend behandelt findet. Man kann kurz folgende Definition aussprechen: Es sei jedem endlichen System von natürlichen Zahlen $n_1, n_2 \dots n_k$ ein Ding $A_{n_1 n_2 \dots n_k}$ zugeordnet, wobei unter den zugeordneten Dingen speziell ein gewisses Ding N , welches wir Null Ding nennen wollen, auftreten kann. Bildet man nun für alle jene Folgen von natürlichen Zahlen $n_1, n_2 \dots n_k, \dots$, für welche für kein natürliches k das Ding $A_{n_1 \dots n_k}$ mit N identisch ist, die Folge der Dinge

$$A_{n_1}, A_{n_1 n_2}, \dots, A_{n_1 n_2 \dots n_k}, \dots,$$

so heißt die Menge aller dieser Folgen von Dingen, die durch das Dingsystem $A_{n_1 n_2 \dots n_k}$ erzeugte Verzweigungsmenge. Von diesem Begriff aus gelangt man in wenigen Worten einerseits zur Brouwerschen Mengendefinition, andererseits zum Begriff der analytischen Menge. Leider wird auch der letztere im Fraenkelschen Buche nicht auseinandergesetzt, obwohl er, ganz abgesehen von seiner mathematischen Bedeutung, auch im Zusammenhange mit Grundlagenfragen erwähnenswert wäre mit Rücksicht auf seine engen Beziehungen zu den Bemühungen der großen Funktionentheoretiker der Pariser Schule um einen konstruktiven Mengenbegriff.

Aus der Tatsache der Äquivalenz der ganz auf Cantorschem Boden stehenden Definition der analytischen Menge mit der Brouwerschen Mengendefinition ergibt sich, wenn man noch bedenkt, daß das Kontinuum der reellen Zahlen eine analytische Menge ist, erstens, daß die Differenzen zwischen der Brouwerschen und der üblichen Behandlung der Mengenlehre erheblich geringer sind, als Fraenkel an mehreren Stellen (z. B. S. 222—226, 239) annimmt, und zweitens, daß gewisse Bemerkungen Fraenkels einer Präzisierung fähig sind. Wenn der Verfasser z. B. erwähnt, daß der Ausdruck „Spezies“ bei Brouwer *etwa* dem üblichen „Menge“ entspricht (S. 237), so kann man genau sagen, daß die Brouwerschen Spezies mit den beliebigen Teilmengen analytischer Mengen identisch sind. Wenn (S. 226) erwähnt wird, daß Borels Standpunkt hinsichtlich der Überabzählbarunendlichen über den Brouwerschen hinausgeht, so kann dies dahin präzisiert werden, daß Brouwers Mengen- und Speziesdefinitionen letzten Endes auf die Betrachtung von Mengen beliebiger Folgen natürlicher Zahlen hinauslaufen, während Borel nur gesetzmäßig festgelegte Folgen natürlicher Zahlen als Elemente wohldefinierter Mengen zulassen will.

Sehr berechtigt erscheint Fraenkels Frage (S. 228 und 382) nach präzisen Angaben über den Begriff der Konstruktivität und Konstruktion. Fraenkel läßt es dahingestellt, „ob man den Begriff der Konstruktion selbst erschöpfend zergliedern kann“. Nach Ansicht des Referenten ist dies höchstens auf Grund willkürlicher Festsetzungen und sicher nicht eindeutig möglich. Kein mathematischer Begriff bietet sich eindeutig und zwangsläufig dar und ein Begriff über das mathematische Schließen erst recht nicht. Daß hinsichtlich dieser letzteren besonders komplizierten Begriffe derartige Fragen überhaupt auftreten konnten, rührt nur daher, daß Erkenntnistheoretiker und Methodiker vor der Anwendung ihrer allgemeinen Ergebnisse auf ihre eigenen Wissenschaften oftmals zurückschrecken.

Nach dem Intuitionismus werden die nicht-prädikativen Begriffsbildungen behandelt. Stark erweitert sind in der neuen Auflage die Ausführungen über Russells Logizismus. Diesem System gibt Fraenkel am Ende des Buches in einem persönlichen Glaubensbekenntni den Vorzug vor den übrigen Versuchen zur Begründung der Mengenlehre (ohne übrigens Russells empiristischen Standpunkt in der Erkenntnistheorie zu teilen). Daß Fraenkel in der neuen Auflage auch zu den Russellschen Gedanken, die in Deutschland bisher entschieden zu wenig bekannt waren, einen leichtfaßlichen Zugang öffnet, ist sicherlich lebhaft zu begrüßen.

Das fünfte letzte Kapitel ist der Axiomatik gewidmet. Eingehend werden Zermelos Aussonderungsaxiom, das von Fraenkel stammende präzisierte Aussonderungsaxiom, das Potenzmengenaxiom und das Auswahlaxiom behandelt.

Nach einer ausführlichen Würdigung der axiomatischen Methode folgt eine Besprechung der Diskussion zwischen Hilbert und Brouwer. Mit Recht bemerkt Fraenkel, daß beide Teile (im Gegensatz zu Russell) wenigstens in gewissem Maße die vollständige Induktion zum Ausgang nehmen und erst in ihrer Verwendung die Differenzen einsetzen. Hilbert verwendet sie, um für die wichtigsten Teile der Mathematik, insbesondere für die klassische Theorie des Zahlenkontinuums, die Widerspruchsfreiheit zu erweisen. Brouwer gebraucht sie bloß zu mathematischen Konstruktionen, rechnet nur diese zur Mathematik und erklärt metamathematische Betrachtungen, insbesondere Widerspruchsfreiheitsbeweise, für mathematisch bedeutungslos. Um allerdings eine Theorie des Zahlenkontinuums zu entwickeln, nimmt Brouwer durch seine Mengen- und Speziesdefinition Konstruktionen von einem Umfang vor, der den der Hilbertschen metamathematischen Konstruktionen erheblich überschreitet, so daß Hilbert diesen Ergebnissen erst vor dem Forum die Widerspruchsfreiheit eine Bedeutung zuerkennen will, zumal eine Präzisierung des Konstruktionsbegriffes, wie erwähnt, aussteht.

Mit einer schönen Würdigung von Cantors unvergänglicher Leistung schließt das Buch. Eines wird hierbei ein Kenner der neuesten Entwicklung der Cantorschen Ideen vermissen. Wir meinen natürlich nicht etwa eine Theorie der Punktmengen; denn eine auch nur einigermaßen ausführliche Behandlung

dieses Wissenszweiges, der dank den neuesten Errungenschaften die abstrakte Mengenlehre an Umfang weit übertrifft, wäre ja mit der ganzen, so wohlbewährten Anlage des auf Kardinalzahl-, Ordinalzahl- und Grundlagentheorie zugeschnittenen Fraenkelschen Buches völlig unverträglich; und man versteht daher, daß Fraenkel so wie in früheren Auflagen nur die primitivsten Beispiele von Punktmengen als Anwendung der Ordnungslehre anführt und im übrigen auf die diesbezügliche Literatur verweist. Was wir in dieser Hinsicht vermissen, ist vielmehr bloß ein methodischer Hinweis, nämlich auf die noch durchaus nicht genügend bekannte Tatsache der Unabhängigkeit der mengentheoretischen Geometrie von den wichtigsten gegen die abstrakte Mengenlehre gerichteten Einwänden. Es wäre eine dankbare Aufgabe, eine Axiomatik der Punktmengenlehre aufzustellen, welche die und nur die in den fruchtbaren Teilen der mengentheoretischen Geometrie verwendeten Schlüsse festlegt. Da käme vor allem ein klassenkalkülartiger Unterbau in Betracht, denn in der Geometrie handelt es sich stets um die Teilmengen eines gegebenen Raumes. Sodann wäre das Aussonderungsaxiom zu erörtern, welches aber zur Bildung von „geometrischen Orten“ und „Gesamtheiten“ auch in der elementaren und analytischen Geometrie verwendet wird, und endlich ein ganz spezieller Fall des Auswahlaxioms (nämlich die Auswahl aus einer abzählbaren Folge von Mengen), welcher auch zur Begründung der Lehre von den reellen Zahlen, also in der gesamten Analysis, unentbehrlich ist. Keine Rolle spielen in den fruchtbarsten Teilen der Punktmengenlehre dagegen das Potenzmengenaxiom, höhere Mächtigkeiten als die des Kontinuums oder gar derart unfruchtbare, ja verdächtige Begriffsbildungen, wie jene in der abstrakten Mengenlehre bei der Einführung des Unendlichen betrachtete Menge, welche neben jedem Element auch die aus diesem Element bestehende Menge als Element enthält. Typenvermischungen, wie die Betrachtung von Mengen, welche teils Elemente, teils Mengen von diesen Elementen enthalten, kommen in der mengentheoretischen Geometrie überhaupt nie vor (geschweige denn iterierte Typenvermischungen), und es besteht in ihr auch nirgends der geringste Anlaß zu solchen Mengenbildungen.

Derartige Ausführungen wären in einer Einleitung in die Mengenlehre vom methodischen Gesichtspunkt wie auch zur Würdigung Cantors wohl von Interesse. Zeigen sie doch, daß von der teilweise kontroversen abstrakten Mengenlehre ganz abgesehen, die Geometrie Cantor einen Fortschritt verdankt, dem an Bedeutung höchstens die Entdeckung der analytischen Methoden an die Seite gestellt werden kann und dessen methodische Verbürgtheit jener der gesamten Analysis nicht nachsteht! Aber auch so ist zu wünschen und zu erwarten, daß Fraenkels ausgezeichnetes Lehrbuch der abstrakten Mengenlehre und seine anregenden und reichhaltigen Ausführungen über Grundlagenfragen den gewaltigen Ideen Cantors weitere Verbreitung und neue Bewunderer verschaffen werden.

Menger.

B. Russell, Unser Wissen von der Außenwelt. Übersetzt von W. Rothstock. Bei F. Meiner, Leipzig 1926. VIII + 331 Seiten. RM 10,—, geb. RM 12,—.

„In der gegenwärtigen Philosophie lassen sich drei Hauptrichtungen erkennen . . . Die erste von ihnen, die ich die klassische Tradition nennen möchte, leitet sich in der Hauptsache von Kant und Hegel her und stellt einen Versuch dar, Methoden und Ergebnisse der großen, spekulativen Denker seit Plato den Bedürfnissen der Jetztzeit anzupassen. Die zweite Richtung, die wir als Evolutionismus bezeichnen können, kam durch Darwin zu ihrer großen Bedeutung, während als ihr erster Vertreter auf philosophischem Gebiete Herbert Spencer angesprochen werden muß. In jüngster Zeit zeigt sich diese Richtung, besonders unter dem Einfluß von William James und Henri Bergson, viel kühner und mehr zu durchgreifenden Neuerungen geneigt, als sie zur Zeit Herbert Spencers gewesen ist. Die dritte Richtung wollen wir in Ermangelung einer besseren Bezeichnung „logischen Atomismus“ nennen. Ihre Methode ist in Anlehnung an die kritischen Untersuchungen der Mathematiker langsam entstanden. Diese Art des Philosophierens, die einzige von allen, die ich glaube vertreten zu können, hat bis jetzt noch nicht viele über-

zeugte Anhänger, aber der von Harvard ausgehende „Neorealismus“ ist sehr stark von demselben Geiste durchdrungen. Meines Erachtens liegt hier ein ähnlicher Fortschritt vor, wie er durch Galilei in der Physik hervorgerufen wurde: beweisbare Einzelergebnisse treten an die Stelle unbewiesener, auf das Ganze gehender Behauptungen, für die man sich nur auf die Einbildungskraft berufen kann.“ Eine kritische Durchmusterung unserer „common knowledge“ zeigt, daß als „harte Daten“, die normalerweise nicht bezweifelt werden können, nur die Sinneswahrnehmungen und die Logik (die moderne, nicht die scholastische Logik) übrig bleiben; auf diesen harten Daten ist nun alles aufzubauen, alle in der Wissenschaft geläufigen Begriffe sind, sofern sie Geltung behalten sollen, durch logische Konstruktion auf die Daten der Wahrnehmung zurückzuführen; die „Dinge“ werden als die Klassen ihrer Erscheinungen definiert, die Raumpunkte als gewisse Klassen von Körpern, die Zeitmomente als gewisse Klassen von Ereignissen; eine metaphysische Existenz wird ihnen in Befolgung der berühmten Devise Occams nicht zugesprochen. Die an den Begriff der Bewegung sich knüpfenden Probleme geben Anlaß zur Besprechung des Stetigkeitsbegriffes der Mathematik und dieser wieder gibt Anlaß zur Besprechung der Lehre vom Unendlichen im Sinne der Mengenlehre. In einem Schlußkapitel wendet Russell seine Methoden noch auf Kausalität, Ursachbegriff und das Problem des freien Willens an. — So verlockend es wäre, auf Einzelheiten einzugehen, wir müssen es uns versagen, da wir nicht für eine philosophische, sondern für eine mathematische Zeitschrift schreiben. Wir müssen uns darauf beschränken, dieses Werk des großen englischen Philosophen den Mathematikern und Physikern angelegentlichst zu empfehlen, als Musterbeispiel dafür, wie die von der Mathematik herkommenden Methoden berufen sind, die Philosophie zu befruchten und zu reformieren. *Hans Hahn.*

O. Becker, Mathematische Existenz. Untersuchungen zur Logik und Ontologie mathematischer Phänomene (Sonderabdruck aus „Jahrbuch für Philosophie und phänomenologische Forschung“, Band VIII). Bei Niemeyer Halle an der Saale, 1927. VII und 369 Seiten.

Der Intuitionismus Brouwers und der Formalismus Hilberts werden einer sehr ausführlichen philosophischen, oder, sagen wir lieber phänomenologischen, an Husserl und noch mehr an Heidegger orientierten Betrachtung unterzogen, bei der Brouwer den Sieg davonträgt, da Hilberts Formalismus dem „phänomenologischen Grundprinzip der Ausweisbarkeit“ nicht standhalte („Dieses Prinzip des Zugangs, der möglichen originären Gegebenheit, das auch als das ‚Prinzip des transzendenten Idealismus‘ bezeichnet zu werden pflegt, besteht in der Forderung, daß jedes echte Phänomen seinem eigenen Seinssinn nach sich ausweisen müsse, für den, dem seine adäquate Erfassung gelingt, in einem Akt originärer Anschauung“); im Formalismus drehe es sich nur um Konsequenz, im Intuitionismus um Wahrheit, die intuitionistische Mathematik gebe Beweise, die formalistische nur Ableitungen, darum sollte Hilbert nicht von einer „Beweistheorie“, sondern von einer „Ableitungstheorie“ sprechen. In sehr lesenswerten, historischen Ausführungen wird dargetan, daß Hilbert der Fortsetzer von Plato und Leibniz, Brouwer der Fortsetzer von Aristoteles und Kant ist. Immer gestützt auf das Prinzip des Zugangs, leugnet der Verfasser das aktual Unendliche, erkennt aber überraschenderweise Cantors transfinite Ordinalzahlen an, da zu diesen die (in einer Erzählung von Dostojewsky geschilderte) iterierte „Reflexion auf sich selbst“ den erforderlichen Zugang liefere, der zurückführt zur ursprünglichen Einführung dieser Zahlen durch Cantors Erzeugungsprinzip, während der später von Cantor eingeschlagene Weg, der die Theorie der wohlgeordneten Mengen voranstellt, für den Verfasser, der ja das zuganglose aktual Unendliche leugnet, ungangbar ist. Auch verhält sich der Verfasser zustimmend zu Hilberts Ansatz eines Beweises für Cantors Vermutung, daß das Kontinuum die Mächtigkeit der zweiten Zahlklasse hat, nur daß seine Interpretation dieses Beweisansatzes eine andere ist. Er will die ganze Überlegung aus der formal-mathematischen Sphäre in die meta-mathematische überstellen. Im ganzen wird der Mathematiker viele Einwendungen gegen Beckers Überlegungen vorzubringen haben, aber auch viel Anregung

und manche Belehrung aus diesem Buche schöpfen können. Es muß anerkennend festgestellt werden, daß — ganz anders als dies früher vielfach üblich war — ein Philosoph sich ernstlich um das Verständnis der mathematischen Theorien bemüht hat, bevor er es unternahm, darüber zu reden und zu schreiben.

Hans Hahn.

K. Menger, Dimensionstheorie. Bei B. G. Teubner, Leipzig u. Berlin 1928. IV + 316 Seiten. Geh. RM 22,—, geb. RM 24,—.

Nur wenige Jahre sind verstrichen seit dem Erscheinen der grundlegenden Arbeiten von Menger und Urysohn zur Dimensionstheorie; aber diese Theorie hat heute schon einen solchen Umfang und eine solche Abrundung erreicht, daß eine zusammenfassende Darstellung in einem Lehrbuche als durchaus zeitgemäß, ja als ein Gebot der Stunde erscheint. Und da das vorliegende Lehrbuch keine Spezialkenntnisse voraussetzt, vielmehr alle erforderlichen Hilfsmittel aus der Theorie der Punktmengen selbst entwickelt, bietet es jedem Mathematiker einen bequemen Zugang zu den Resultaten der Dimensionstheorie, deren Bedeutung und Interesse ja stark über das Spezialgebiet der Punktmengenlehre hinausgreift. Die Darstellung beginnt mit einer Einführung der separablen Räume, die sich engstens an die Einführung der reellen Zahlen von der abzählbaren Menge der rationalen Intervalle her anschließt. Es werden die Abgeschlossenheitseigenschaften der Mengen entwickelt; eine eingehende Darstellung finden die Sätze über die Begrenzung offener Mengen und die Überdeckungssätze, die für die weiteren Untersuchungen grundlegend sind; hier ist von besonderem Interesse das von Menger formulierte Überdeckungstheorem für halbkompakte Mengen (S. 46), zu dessen tieferem Verständnis Hurewicz wesentlich beigetragen hat. Nach einer kurzen Erörterung der Dichteitseigenschaften wird mit Urysohn die Einbettbarkeit eines beliebigen separablen Raumes in den abzählbardimensionalen Grundquader bewiesen, wodurch auch das Problem der Metrisierung miterledigt ist; aus dem Inhalte des ersten einleitenden Kapitels sei noch hervorgehoben ein sehr eleganter (von Kuratowski und Hurewicz herrührender) Beweis des Brouwerschen Theorems von der Existenz abgeschlossener Mengen, die hinsichtlich einer Eigenschaft E irreduzibel sind. Kapitel II bringt die Definition des Dimensionsbegriffes, Kapitel III behandelt den Summensatz (die Summe abzählbar vieler höchstens n -dimensionaler abgeschlossener Mengen ist höchstens n -dimensional) und seine Verallgemeinerungen; ferner den Zerspaltungssatz (eine n -dimensionale Menge ist Summe von $n+1$, aber nicht von weniger als $n+1$ nulldimensionalen Mengen) und den Trennbarkeitssatz (damit ein Raum höchstens n -dimensional sei, ist notwendig und hinreichend, daß je zwei fremde abgeschlossene Teilmengen durch eine höchstens $(n-1)$ -dimensionale Menge trennbar seien). Es folgt die Theorie der rational n -dimensionalen Räume (die Summe von $n+1$ nulldimensionalen Mengen sind, von denen eine abzählbar ist) und Hurewicz's Theorie der Normalbereiche. Kapitel IV handelt von der dimensionellen Raumstruktur; die wichtigsten Sätze dieser Art besagen, daß in einem kompakten Raume R die Menge R^k aller Punkte, in denen R mindestens k -dimensional ist, in jedem ihrer Punkte mindestens k -dimensional ist, während in einem beliebigen separablen Raum von ihr nur gesagt werden kann, daß sie mindestens $(k-1)$ -dimensional ist; doch ist die abgeschlossene Hülle von R^k in jedem Punkte von R^k mindestens k -dimensional. Kapitel V ist dem Zerlegungssatze gewidmet, der besagt, daß ein n -dimensionaler Raum Summe von endlich vielen beliebig kleinen Stücken ist, die zu je k einen höchstens $(n-k+1)$ -dimensionalen Durchschnitt haben ($k=2, 3, \dots, n+2$), während bei jeder Zerlegung in endlich viele hinlänglich kleine abgeschlossene Mengen notwendig Mengen- k -tupel mit mindestens $(n-k+1)$ -dimensionalem Durchschnitt auftreten ($k=2, 3, \dots, n+1$). Kapitel VI (Die Zusammenhangseigenschaften der Räume) behandelt den Zusammenhangsbegriff nach Hausdorff, die verstreuten Mengen nach Sierpiński und seinen Schülern, sodann die erst auf dem Boden der Dimensionstheorie entstandenen Begriffe des höherstufigen Zusammenhanges, der höherstufigen Continua, der n -dimensionalen Cantorschen Mannigfaltigkeit. Kapitel VII handelt von den stetigen Abbildungen; es wird die Theorie der Streckenbilder vorgeführt, dann die

Theorie der dimensionserniedrigenden und der dimensionserhöhenden stetigen Abbildungen, endlich wird die Invarianz des Dimensionsbegriffes gegenüber topologischen Abbildungen bewiesen. Kapitel VIII (Die Dimensionsverhältnisse in Cartesischen Räumen) bringt den Nachweis, daß der R_n , der n -dimensionale Raum der analytischen Geometrie, auch n -dimensional im Sinne der Dimensionstheorie ist; genauer: die n -dimensionalen Punktmengen des R_n sind identisch mit jenen seiner Punktmengen, die einen offenen Teil enthalten. Kapitel IX bringt den Nachweis der höchst überraschenden, von Menger gefundenen Tatsache, daß jeder endlich-dimensionale separable Raum mit einer Teilmenge eines Cartesischen Raumes R_n homöomorph ist, daß insbesondere jeder n -dimensionale separable Raum mit einer Teilmenge des R_{2n+1} homöomorph ist. Ein Schlußkapitel gibt eine übersichtliche Zusammenstellung der erreichten Hauptergebnisse und Angaben über noch ungelöste Probleme. Besonders hervorgehoben seien noch die am Schlusse der einzelnen Paragraphen angebrachten, oft recht ausführlichen historischen Noten, auf Grund derer sich jedermann ein eigenes Urteil über die Entstehung und Autorschaft der einzelnen Teile der Dimensionstheorie bilden kann, die ja — wie bekannt — gleichzeitig und unabhängig in ihren wesentlichen Zügen von Menger und dem auf tragische Weise in jungen Jahren der Forschung entrisenen Urysohn aufgebaut wurde. — Es ist kein Zweifel, daß wir es hier mit einem der interessantesten Bücher der mathematischen Literatur der letzten Jahre zu tun haben, das in anregender und nicht schwieriger Form in ein völlig neues und überaus wichtiges Kapitel der Mathematik einführt.

Hans Hahn.

C. Carathéodory, Vorlesungen über reelle Funktionen. Zweite Auflage. Bei B. G. Teubner, Leipzig u. Berlin 1927. X + 718 Seiten. RM 29,—.

Die zweite Auflage dieses ausgezeichneten Werkes wurde durch mechanischen Neudruck hergestellt, so daß es nicht möglich war, größere Änderungen vorzunehmen. Die wichtigsten Änderungen im Texte betreffen den Abschnitt über Punktmengen von nicht meßbarem Inhalt (§ 332—334), der im Anschlusse an eine Comptes-Rendus-Note von J. Wolff umgearbeitet wurde, den Abschnitt über die Darstellung halbstetiger Funktionen als Grenzwerte stetiger Funktion (§ 367), der im Anschlusse an eine bekannte Abhandlung von F. Hausdorff umgearbeitet wurde, und den Abschnitt über die Erweiterung stetiger Funktionen (§§ 541—543), der — wieder im Anschlusse an Hausdorff — auf Grund des vom Referenten bewiesenen Einschließungssatzes neu bearbeitet wurde. Am Schlusse des Bandes findet man als Note I ein von H. Bohr herrührendes Beispiel, das zeigt, daß die behandelte, im Vitalischen Überdeckungssatze auftretende Regularitätsbedingung für die Gültigkeit dieses Satzes wesentlich ist, ferner als Note II den Nachweis, daß, wenn $\mu_a(A)$ und $\mu_i(A)$ das äußere bzw. innere Maß einer regulären Maßfunktion bedeuten, das arithmetische Mittel $\frac{1}{2}(\mu_a[A] + \mu_i[A])$ wieder eine Maßfunktion ist, die aber im allgemeinen nicht regulär ist, wodurch man besonders einfache Beispiele für nichtreguläre Maßfunktionen erhält. Es folgt dann noch ein sehr übersichtliches, 105 Nummern enthaltendes Literaturverzeichnis. — Zweifelloos wird auch die zweite Auflage dieses Werkes, das so viel dazu beigetragen hat, die an Lebesgue anknüpfende Entwicklung der mathematischen Forschung weiteren Kreisen von Mathematikern zugänglich zu machen, den ihm gebührenden Erfolg finden.

Hans Hahn.

G. Kowalewski, Grundzüge der Differential- und Integralrechnung. Vierte, verbesserte Auflage. Bei Teubner, Leipzig u. Berlin 1928. V + 417 Seiten. RM 16,—.

Eine neue Auflage dieses ausgezeichneten Lehrbuches der Differential- und Integralrechnung, dessen Vorzug es ist, Exaktheit und Lesbarkeit miteinander zu verbinden. In einem ersten Anhang sind die einfachsten Sätze über Determinanten behandelt, ein zweiter, in dieser Auflage neu hinzugekommener Anhang behandelt die Fredholm'schen Determinanten und Minoren durch Grenzübergang vom Endlichen aus und deutet ihre Verwendung bei Auflösung der linearen Integralgleichungen zweiter Art an.

Hans Hahn.

L. Bieberbach, Differential- und Integralrechnung. Band II: Integralrechnung (Teubners Mathematische Leitfäden, Band 5). Dritte, vermehrte und verbesserte Auflage. Bei Teubner, Leipzig u. Berlin 1928. VI + 150 Seiten. RM 5,80.

Mit aufrichtiger Freude sei das Erscheinen einer neuen Auflage dieses ausgezeichneten Lehrbuches angezeigt, das bei knappem Umfange und mäßigem Preise dem Studierenden nicht nur alles für ihn nötige bringt, sondern es ihm durch die lebendige Form der Darstellung auch wirklich nahebringt und das daher in jeder für Anfänger bestimmten Vorlesung den Hörern empfohlen werden sollte.

Hans Hahn.

C. G. J. Jacobi, Theorie der elliptischen Funktionen. Herausgegeben von Adolf Kneser. (Ostwalds Klassiker Nr. 224). Leipzig akad. Verlagsgesellschaft m. b. H., 1927. RM 3,20.

Ein Ausgabe der bekannten Vorlesung Jacobis nach der Nachschrift Borchardts von 1838, mit wenigen Anmerkungen des Herausgebers, die sich auf das Umkehrproblem und die lineare Transformation beziehen und einen Anhang über die Landensche Transformation.

Wirtinger.

C. F. Gauss. Bestimmung der Anziehung eines elliptischen Ringes. Nachlaß zur Theorie des arithmetisch-geometrischen Mittels und der Modulfunktion. Übersetzt und herausgegeben von H. Geppert. (Sammlung Ostwald, 225). Akad. Verlagsgesellschaft, Leipzig 1927. RM 11,60.

In diesem Bändchen sind außer der zuerst genannten auch die übrigen erst später bekanntgewordenen und noch später verstandenen Untersuchungen von Gauss zur Theorie der elliptischen Funktionen zum erstenmal in deutscher Übersetzung mit den nötigen Erläuterungen bequem zugänglich gemacht.

Wirtinger.

L. Euler, Drei Abhandlungen über die Auflösung der Gleichungen. Übersetzt und herausgegeben von S. Breuer. (Ostwalds Klassiker Nr. 226.) Bei Akad. Verlagsgesellschaft, Leipzig 1928. RM 5,20.

I. Eine Vermutung über die Wurzelformen der Gleichungen aller Grade. Euler vermutet, daß zu jeder Gleichung

$$(1) \quad x^n = a x^{n-2} + b x^{n-3} + \dots$$

eine Resolvente existiert

$$(2) \quad z^{n-1} = \alpha z^{n-2} - \beta z^{n-3} + \gamma z^{n-4} \dots$$

mit den $n-1$ Wurzeln $A, B, C \dots$ derart, daß die n Wurzeln von (1) unter den n^{n-1} Werten des Ausdrucks

$$(3) \quad x = \sqrt[n]{A} + \sqrt[n]{B} + \sqrt[n]{C} + \dots$$

enthalten seien. Da ihm dies allgemein nicht gelingt, so betrachtet er zunächst spezielle Resolventen.

II. Von der Auflösung der Gleichungen aller Grade. Euler versucht die Gleichung zu finden, die als Wurzelform den n -wertigen Ausdruck hat:

$$(4) \quad x = \mathfrak{A} \sqrt[n]{v} + \mathfrak{B} \sqrt[n]{v^2} + \dots \mathfrak{D} \sqrt[n]{v^{n-1}},$$

um dann umgekehrt die allgemeine Gleichung n ten Grades zu lösen.

III. Zahllose Formen von Gleichungen aller Grade, deren Auflösung sich um dann umgekehrt die allgemeine Gleichung n ten Grades zu lösen. angeben läßt. Euler behandelt in diesem Abschnitt spezielle auflösbare Gleichungen.

gen. Herr Breuer hat die drei genannten Abhandlungen in klares Deutsch übersetzt und mit sehr vielen Anmerkungen versehen. Diese bringen eine Würdigung der Abhandlungen, eine kurze Übersicht, ferner Erklärungen und Ergänzungen, die alle von modernem Standpunkte aus gegeben werden. *Hofreiter.*

J. Horn, Gewöhnliche Differentialgleichungen. (Göschens Lehrbücherei, Band 10.) Bei Gruyter & Co., Berlin und Leipzig 1927. RM 9,—, geb. RM 10,50.

Das Buch ist eine sehr gute Einführung in die mehr elementaren Teile dieser Theorie. Es behandelt der Reihe nach die elementaren Integrationsmethoden, Existenzbeweise, schrittweise Annäherung, numerische und graphische Näherungsmethoden. Sodann ausführlicher die linearen Differentialgleichungen, die Abhängigkeit der Lösungen von Parametern und Anfangswerten. Als Anwendung die Aufsuchung periodischer Lösungen und endlich die einfachsten Singularitäten nichtlinearer Differentialgleichungen wird durch zweckmäßige Beispiele erläutert und manche Ausführung dem Leser überlassen, wenn sie nach einem bereits durchgeführten Muster geschehen kann. Dadurch wird Raum gespart und der Leser nicht ermüdet. Das Buch kann also gewiß zum Studium und als Ratgeber bei Anwendungen empfohlen werden.

Bei der Gelegenheit sei aber dem Berichterstatter eine allgemeine Bemerkung getattet, zu der er durch Anfragen im Verlaufe seiner Lehrtätigkeit veranlaßt wird. Die Erklärung einer gewöhnlichen Differentialgleichung als einer Gleichung zwischen einer unabhängigen Veränderlichen, einer Funktion y von x und einem oder mehreren Differentialquotienten $y', y'', y \dots^{(n)}$ bedarf einer Ergänzung.

Es muß nämlich die höchste Derivierte eine gegebene Funktion der Derivierten niedrigerer Ordnung und der beiden Variablen sein. Die Bedeutung des Wortes gegeben wird klar, wenn man etwa die Aufgabe stellt, alle Kurven zu bestimmen, welche durch eine vorgegebene Berührungstransformation in sich übergehen. Sei $y = y(x)$ die gesuchte Kurve und $\xi = \varphi(x, y, p)$, $\eta = \psi(x, y, p)$, $\zeta = \chi(x, y, p)$ die Gleichungen der gegebenen Berührungstransformation, so ist die von der gesuchten Kurve zu erfüllende Gleichung

$$\psi(x, y, y') = y[\varphi(x, y, y')].$$

Aber diese hat einen ganz anderen Charakter, wie eine Differentialgleichung im gewöhnlichen Sinn und man überzeugt sich an einfachen Beispielen leicht, daß dabei eine abzählbare Menge willkürlicher Funktionen auftreten kann. Es tritt eben hier die Funktion $y(x)$ selbst auf, welche nicht gegeben ist. Der Berichterstatter will keineswegs einen Tadel damit aussprechen, sondern nur auf diesen, wie es scheint, bisher wenig beachteten Umstand hinweisen. *Wirtinger.*

E. Wicke, Konforme Abbildungen. (Mathematisch-Physikalische Bibliothek, herausgegeben von W. Lietzmann und A. Witting, Band 73.) Bei Teubner, Leipzig 1927. 59 Seiten. RM 1,20.

Entsprechend der Absicht, mit der die oben genannte Sammlung herausgegeben wird, ist die vorliegende kleine Schrift ein Versuch, eine einfache, nur etwa die Kenntnisse eines absolvierten Mittelschülers voraussetzende Einführung in das Wesen der konformen Abbildung zu geben. Zunächst wird die Inversion besprochen. Dann folgt ganz kurz eine Erörterung des Begriffes der komplexen Zahl und der analytischen Funktion mit der durch sie vermittelten konformen Abbildung. Genauer ausgeführt werden die Fälle der allgemeinen Kreisverwandtschaft, wie sie durch die linear gebrochene Funktion geliefert wird, ferner die durch $Z = z^2$ und einige andere quadratische Funktionen vermittelten Abbildungen, wobei sich Gelegenheit zur Einführung des Begriffes der Riemannschen Fläche ergibt. Nach einigen hydrodynamischen Anwendungen wird zum Schluß die Funktion e^z und die durch sie gelieferte Abbildung besprochen. Wenn vielleicht auch an einzelnen Stellen der logische Aufbau klarer hervortreten könnte, so ist doch das Buch mit großem pädagogischem Geschick ge-

schrieben und könnte Studierenden der ersten Semester als Vorschule zur Theorie der konformen Abbildungen dienen. *Helly.*

K. Weierstrass, Mathematische Werke; Band VII. Vorlesungen über Variationsrechnung. Bearbeitet von R. Rothe. Akademische Verlagsgesellschaft, Leipzig 1927. VIII + 328 Seiten. RM 45,—.

Endlich die langerwartete authentische Ausgabe der berühmten Vorlesungen von Weierstrass über Variationsrechnung. Die Motive für die an sich unverständliche Verzögerung in der Herausgabe dieser Vorlesungen, die zu verschiedenen Gerüchten Anlaß gegeben hatte, finden sich im Vorworte dunkel angedeutet: „In dem ursprünglichen Plane der Weierstrassischen Werke war die Herausgabe eines Bandes, der die Vorlesungen über Variationsrechnung enthielt, nicht vorgesehen, aus Gründen, die sich jetzt nicht mehr genau feststellen lassen, vermutlich, weil eine gesonderte Veröffentlichung dieser Vorlesung von anderer Seite beabsichtigt war. Erst als eine Verwirklichung dieser Absicht nicht mehr in Frage kam, wurde ein Band über Variationsrechnung in den Gesamtplan einbezogen.“ Einer Veröffentlichung vor 25 Jahren wäre höchst aktuelles Interesse zugekommen, heute ist dieses Interesse nur mehr ein historisches, aber das historische Interesse ist auch heute noch ein außerordentliches. Der Veröffentlichung liegen zugrunde: eine von H. Burckhardt angefertigte, von H. A. Schwarz abgeschriebene Ausarbeitung der Vorlesung von 1882, ferner Ausarbeitungen der Vorlesungen von 1875 und 1879. Sämtliche Korrekturbogen wurden von Bieberbach und Carathéodory einer kritischen Durchsicht unterzogen. Der Darstellung der Variationsrechnung ist die Lehre von der Maxima und Minima vorausgeschickt. Zu ausführlicher Darstellung gelangen das einfachste Problem der Variationsrechnung und das einfachste isoperimetrische Problem (beide in Parameterdarstellung), in dem Umfange, der durch die heute klassischen und jedem, der sich mit Variationsrechnung befaßt hat, wohl bekannten Weierstrassschen Forschungen abgesteckt ist. An speziellen Problemen werden ausführlich behandelt: die kleinste Rotationsfläche, die Brachistochrone, die geodätischen Linien, der Rotationskörper kleinsten Widerstandes, das isoperimetrische Problem im engeren Sinne, die Gestalt eines schweren hängenden Fadens. *Hans Hahn.*

M. Lindow, Numerische Infinitesimalrechnung. Bei Dümmler, Berlin 1928. RM 15,—.

Es ist bekannt, daß gerade die Astronomie besonders viel mit Zahlenrechnungen zu tun hat, und hierin eine Tradition von ehrfurchtgebietender Dauer aufzuweisen hat. Es ist daher stets hochinteressant, eine Darstellung der Interpolationsrechnung aus der Feder eines Astronomen, wie sie hier geboten wird kennen zu lernen. In der Tat kann nicht nur der angehende Astronom, sondern auch jeder, der überhaupt sich mit der angewandten Mathematik befaßt, aus dem Buche sehr viel lernen.

Im ersten Kapitel wird die Interpolation durch Polynome und die Interpolation periodischer Funktionen ausführlich besprochen die beiden folgenden enthalten die numerische Differentiation und Integration, das letzte endlich ist der numerischen Lösung von Differentialgleichungen gewidmet, wobei insbesondere das Dreikörperproblem behandelt wird. Überall sorgen ziemlich weit ausgeführte Beispiele dafür, daß der Leser bis zur wirklichen Anwendung der vorgetragenen Regeln vorzudringen vermag.

In den theoretischen Erläuterungen ist die Ausdrucksweise hie und da vom Standpunkt des Mathematikers nicht ganz einwandfrei. Es ist zwar kaum anzunehmen, daß das Studium des Buches hierdurch wesentlich beeinträchtigt werden wird, indessen wäre es doch wünschenswert, daß bei einer neuen Auflage diese wenigen Stellen verbessert würden, so daß das wertvolle Buch auch von diesen mehr äußerlichen Nachteilen frei würde. *Schrutka.*

H. v. Sanden, Mathematisches Praktikum. I. Teil. Mit 17 Figuren und 20 Zahlentafeln (Teubners technische Leitfäden, 27). Bei Teubner, Leipzig 1927. RM 6,80.

Das kleine Buch ist der erste Vertreter einer ganz neuen Gattung. Es bringt Zahlenrechnungen in genauer Durchführung mit zahlreichen Anweisungen und Überlegungen rechentechischer Art. Das Buch eignet sich nicht als Nachschlagewerk und auch nicht zum Lesen, es wird vielmehr seinen Nutzen nur dann zeigen, wenn der Student eine größere Anzahl der vorgeführten Aufgaben, womöglich immer wieder selbst den weiteren Weg suchend, durcharbeitet. Damit soll nicht gesagt sein, daß nicht der Eingeweihte gegebenenfalls auch Auskunft über eine bestimmte Art von Aufgaben daraus schöpfen kann.

Die behandelten Gebiete sind: Rechenhilfsmittel, insbesondere der Rechenschieber; der Taylorsche Satz und Näherungsrechnungen; Auflösung von Gleichungen; Ausgleichung; Differentiation, Integration und Interpolation; harmonische Analyse.

Die im Titel genannten Zahlentafeln sind nicht, wie man wohl vermuten könnte, Hilfstafeln für Rechnungen bestimmter Art, sondern übersichtliche Darstellungen des Rechengangs bei einigen umfangreicheren Beispielen. Sie sind in einem eigenen Heft vereinigt, wodurch eine bequeme Vergleichung mit dem erläuternden Text ermöglicht wird.

Man kann nur wünschen, daß das wertvolle Buch überall, wo angewandte Mathematik betrieben wird, recht eifrig benützt werden möge. *Schrutka.*

Hellinger-Toeplitz, Integralgleichungen und Gleichungen mit unendlich vielen Unbekannten. (Sonderausgabe aus der Enzyklopädie der mathematischen Wissenschaften.) Bei Teubner, Leipzig 1928. RM 18,—.

Das Erscheinen dieses Werkes, das einen Sonderabdruck des gleichnamigen Artikels der Enzyklopädie der mathematischen Wissenschaften darstellt, ist außerordentlich zu begrüßen. Es existiert ja eigentlich kein Lehrbuch der Integralgleichungstheorie, das sie in ihrer Gesamtheit darstellte. Und speziell die neueren Methoden konnten nur in den Originalabhandlungen aufgefunden werden. Es ist daher besonders dankenswert, daß die Verfasser, die ja selbst so viel zum Ausbau der Theorie beigetragen haben, in diesem Buche über den Rahmen eines eigentlichen Enzyklopädieartikels hinausgegangen sind und die Gedankengänge der Beweise und in vielen, sonst schwer zugänglichen Fällen auch ausgeführte Beweise gegeben haben, wobei der Kenner des Gebietes bemerken wird, wieviel eigene Arbeit in dieser Wiedergabe liegt.

Die Anordnung des Stoffes ist derart, daß nach einer historischen Einleitung, in der besonders auf die allmählich klar hervortretenden algebraischen Analogien hingewiesen wird, zunächst ein großer Abschnitt folgt, in dem die verschiedenen Lösungsmethoden systematisch dargestellt werden: Die Fredholm'sche Theorie, die Methode der unendlich vielen Veränderlichen und im Anschluß an letztere eine außerordentlich instruktive Übersicht über den heutigen Stand der Theorie der Gleichungen mit unendlich vielen Veränderlichen. Ferner ein Bericht über Integralgleichungen erster Art (Momentenprobleme), über lineare Funktionaloperationen und schließlich über die verwandten nichtlinearen Probleme. Als zweiter Hauptabschnitt des Werkes folgt nun die „Eigenwerttheorie“ mit den Entwicklungssätzen und der allgemeinen Theorie der quadratischen und bilinearen Formen von unendlich vielen Veränderlichen. Speziell in dieser letzteren Theorie, die ja das Hauptarbeitsgebiet der Verfasser ist, findet sich manches Neue oder zum mindesten im Zusammenhang neuartig Dargestellte.

Die Darstellung ist sehr klar und leicht verständlich und legt besonders Wert darauf, die Zusammenhänge überall anschaulich hervortreten zu lassen.

Daß, bei dem ursprünglichen Zweck der Schrift als Enzyklopädieartikel, sehr vollständige Literaturangaben vorhanden sind, muß nicht besonders betont werden.

Alles in allem ist das Werk für jeden, der auch nur etwas Interesse für die behandelten Gebiete hat, kein Buch, das man kritisiert, sondern eines, über das man sich freut. *Helly.*

Fr. A. Willers, *Methoden der praktischen Analysis*. (Göschens Lehrbücherei, I. Gruppe, Bd. 12.) Bei de Gruyter, Berlin 1928. Geb. RM 21,50.

Nach einer längeren Zeit der Zurücksetzung finden nunmehr die Methoden des Zahlenrechnens eine erfreuliche Pflege. Nachdem bereits mehrere Werke erschienen sind, die auf die Schule von Karl Ränge zurückgehen, liegt nun hier eine Darstellung vor, die durch die vorangegangene engere Zusammenarbeit des Verfassers mit Rudolf Rothe ein besonderes Gepräge erhalten hat. In die Darstellung sind auch zahlreiche Veröffentlichungen in Zeitschriften aus der allerletzten Zeit eingefügt worden.

Über den behandelten Stoff geben die Überschriften der sechs Kapitel Auskunft: Das Zahlenrechnen und seine Hilfsmittel, Interpolation, Angenäherte Integration und Differentiation, Praktische Gleichungslehre, Analyse empirischer Funktionen, Angenäherte Integration gewöhnlicher Differentialgleichungen. Den numerischen Methoden ist, ihrer Wichtigkeit entsprechend, überall der meiste Raum gewidmet, doch sind die graphischen Methoden nicht übergangen, da der Verfasser (im Gegensatz etwa zu E. T. Whittaker) ihnen Bedeutung für die Praxis zuspricht. So ist insbesondere auch der Nomographie eine, wenn auch kurze, so doch sehr übersichtliche Darlegung zuteil geworden. Überall wird die richtige Auffassung des Vorgebrachten durch ausführliche Beispiele gefördert, ein großer Vorteil für den Studenten, der das Buch benützt.

Möge diese neue schätzenswerte Darstellung dazu beitragen, das Interesse für die zahlenmäßige Durchführung der Aufgaben der Mathematik weiter zu fördern.

Schrutka.

E. Pascal, *Repertorium der höheren Mathematik*. Erster Band Analysis, zweiter Teilband unter Mitwirkung von A. Barneck in Berlin, G. Doetsch in Stuttgart, F. Engel in Gießen, R. Fricke in Braunschweig, A. Guldberg in Oslo, H. Hahn in Wien, E. Jahnke in Berlin, H. W. E. Jung in Halle, S. Pascal in Neapel; herausgegeben von E. Salkowski in Hannover. 3. Auflage. 1927. Bei Teubner, Berlin und Leipzig. RM 18,—.

Der Band enthält der Reihe nach gewöhnliche Differential- und Differenzgleichungen, partielle und totale Differentialgleichungen, beides von Guldberg. Hier ist bei der Definition der Differentialgleichungen dasselbe zu bemerken, wie bei dem Lehrbuch von Horn. Totale Differentialgleichungen und Differentialformen von E. Pascal, Transformationsgruppen von Guldberg und Engel, Variationsrechnung von H. Hahn, Funktionstheorie von G. Doetsch, Elliptische Funktionen und Integrale von E. Jahnke †, überarbeitet und ergänzt von A. Barneck, Algebraische Funktionen und ihre Integrale von H. W. E. Jung, Thetafunktionen und Abelsche Funktionen von demselben, Automorphe Funktionen und elliptische Modulfunktionen von R. Fricke.

Ein solches Buch soll in erster Linie dem Mathematiker, der auf einem anderen, als dem eigenen Arbeitsgebiet Belehrung sucht, eine Übersicht über Resultate, Methoden und weitere Literatur bieten. Diesem Zweck entsprechen die einzelnen Beiträge im allgemeinen sehr gut. Der Artikel über Funktionentheorie geht dabei auf ganz moderne Untersuchungen von Fatou, Polya, Landau, Julia und andere ein. Zurück treten aber dabei spezielle Funktionen und insbesondere Funktionen von mehr Veränderlichen. Der Artikel über algebraische Funktionen und ihre Integrale ist durchaus auf der arithmetischen Theorie aufgebaut. Ein Versehen, welches den Unkundigen irreführen könnte, ist auf S. 588 geschehen. Dort sind nicht n unabhängige Variable, sondern überhaupt n Variable gemeint und die Exponenten $\nu_1, \nu_2, \dots, \nu_m$ sind Ordnungen der Differentiale, nicht Potenzexponenten. Auf S. 791 (elliptische Funktionen) ist es notwendig, auch den Zusammenhang der $a b c d, a' b' c' d', a'' b'' c'' d''$ anzugeben.

Auf S. 856 könnte wohl die von Hurwitz gefundene Summe der Ordnungszahlen der Weierstrasspunkte, $(p-1)p(p+1)$, angegeben werden, die sich aus der ebenfalls nicht erwähnten Beziehung der fehlenden Ordnungen zu den Ordnungen des Nullwerdens der Differentiale erster Gattung an einer solchen Stelle ergibt.

Ebenso ist Hurwitz nicht erwähnt in den Literaturangaben S. 860, welche sich auf Transformationen eines algebraischen Gebildes in sich beziehen, was um so mehr von Wichtigkeit ist, weil er zu sehr präzisen Resultaten in bezug auf die Form solcher Gleichungen kommt und eine später von Wiman noch herabgedrückte Grenze für die mögliche Anzahl solcher Transformationen angibt.

Ein dritter Teilband, dessen Inhaltsverzeichnis bereits mit abgedruckt ist, steht in naher Aussicht. *Wirtinger.*

E. Netto, Lehrbuch der Kombinatorik. Zweite Auflage erweitert und mit Anmerkungen versehen von Viggo Brun und Th. Skolem. Bei Teubner, Leipzig 1927, RM 14,—.

Diese zweite Auflage ist bezüglich der ersten 13 Kapitel ein unveränderter Abdruck mit Ergänzungen, die in 19 Noten von Skolem am Ende des Buches bestehen. Neu hinzugekommen sind die Kapitel 14 und 15, deren erstes von Brun stammt und von der Verteilungsfunktion handelt, die in sich die verschiedenen elementaren kombinatorischen Operationen einschließt. Von den Formeln werden Anwendungen auf die Zahlentheorie gemacht. Das 15. Kapitel ist ein Nachtrag von Skolem und handelt von Gruppierungen, kombinatorischen Reziprozitäten und Paarsystemen. *T. Rella.*

O. Perron, Algebra. Zwei Bände. Bei de Gruyter, Berlin 1927. Preis geb. RM 11,50 und 9,50.

Wie der Verfasser in der Vorrede sagt, kommt es ihm vor allem darauf an, den Anfänger in die Methoden der modernen Algebra einzuführen. Wohl kann die Algebra als die Theorie der rationalen Funktionen bezeichnet werden, aber ihre Methode unterscheidet sich grundlegend von den Methoden der Funktionentheorie sowie auch die Veränderlichen (oder besser Unbestimmten) der Algebra ganz etwas anderes sind als die Variablen der Funktionentheorie. Im Mittelpunkt der modernen Algebra steht der Begriff des Körpers und der Begriff des Ringes und es werden daher diese Begriffsbildungen gleich in der Einleitung eingeführt und nicht erst, wie das früher zu geschehen pflegte, erst bei der Behandlung der Galoisschen Theorie. Allerdings hat der Verfasser aus didaktischen Gründen verzichtet, mit abstrakten Körpern zu operieren, wohl aber sind die Beweise so gehalten, daß sie für abstrakte unendliche Körper richtig bleiben. Sehr begrüßenswert erscheint es, daß die Theorie der Resultanten und das Bézout'sche Theorem ausführlich behandelt werden, nicht nur für den Fall von zwei Gleichungen mit zwei Unbekannten. Ferner sei darauf hingewiesen, daß die Frage der Existenz der Wurzeln von algebraischen Gleichungen vom Standpunkt der abstrakten Körpertheorie behandelt wird, indem gezeigt wird, daß durch geeignete Erweiterung des Grundkörpers durch Aufnahme von neuen Rechnungssymbolen wieder ein Körper entsteht, in dem die vorgelegte Gleichung eine resp. so viele Wurzeln besitzt, als ihr Grad beträgt. Der Fundamentalsatz der Algebra wird wohl auch bewiesen, aber mit scharfer Betonung, daß es sich dabei eigentlich um einen Satz der Funktionentheorie handelt, dessen Beweis auch mit rein algebraischen Hilfsmitteln nicht geführt werden kann, der allerdings eine abschließende Antwort bezüglich der Gleichungen mit reellen oder komplexen Koeffizienten gibt, nämlich daß der Körper aller komplexen Zahlen algebraisch abgeschlossen ist. Der vorgetragene Beweis ist im wesentlichen der von Gordan.

Der zweite Band, welcher die Theorie der algebraischen Gleichungen behandelt, fußt wohl auf den Resultaten des ersten, kann aber zum großen Teil ohne Kenntnis des ersten Bandes mit Erfolg gelesen werden. Im ersten Kapitel werden wesentlich Fragen der Analysis behandelt, nämlich die numerische Auflösung von Gleichungen. Das zweite Kapitel handelt von der algebraischen Auflösung der Gleichungen bis zum vierten Grad. Im dritten wird die Theorie der Substitutionsgruppen — mit Außerachtlassung der abstrakten Gruppentheorie — behandelt. Das zentrale vierte Kapitel entwickelt die Galoissche Theorie, die wesentlich auf den Sätzen über die Wurzelexistenz aus dem ersten Band (und nicht auf dem Fundamentalsatz der Algebra) aufgebaut ist; das fünfte Kapitel behandelt die Theorie der Gleichung fünften Grades und hier werden einige wenige Sätze aus der Theorie der elliptischen

Funktionen verwendet, die der Formelsammlung von Schwarz entnommen sind.

Die Darstellung zeichnet sich durch Klarheit und Strenge aus und ist dabei doch so breit gehalten, daß sie auch dem jungen Studenten, an den sich das Werk wohl in erster Linie wendet, keine allzu großen Schwierigkeiten bereiten dürfte.

T. Rella.

H. Hasse, Höhere Algebra II, Gleichungen höheren Grades. (Sammlung Göschen). Bei de Gruyter, Berlin 1927. RM 1,50.

Was über die Vorzüge des ersten Bandes dieses Werkes in einer früheren Besprechung gesagt worden ist, gilt auch für diesen. Es ist Hasse gelungen, die Hauptresultate von Steinitz aus der Theorie der abstrakten Körper, die anscheinend noch nicht die ihnen gebührende Verbreitung gefunden haben, einem weiten Leserkreis leicht faßlich darzustellen.

Ausgehend von den Polynomen einer Unbestimmten x mit Koeffizienten aus einem Grundkörper K wird zuerst die eindeutige Zerlegung dieser Polynome in Primelemente (irreduzible Polynome) nachgewiesen, hierauf gezeigt, daß der Restklassenring nach einem Polynom $f(x)$ dann und nur dann ein Körper ist, wenn $f(x)$ irreduzibel ist. Zum Schluß des ersten Abschnittes wird noch kurz auf die Begriffe Primkörper, Charakteristik, vollkommener und unvollkommener Körper eingegangen, obwohl die folgenden Abschnitte sich auf die Behandlung von vollkommenen Körpern mit unendlich vielen Elementen beschränken. Im zweiten Abschnitt wird der einfache Zusammenhang zwischen den Wurzeln einer Gleichung $f(x) = 0$ in einem existierend angenommenen Erweiterungskörper mit dem Vorhandensein von Linearfaktoren von $f(x)$ in diesem Erweiterungskörper nachgewiesen und darauf fußend die Definition der mehrfachen Wurzeln gegeben, wobei sich für vollkommene Körper der Satz ergibt, daß ein irreduzibles Polynom keine mehrfachen Wurzeln besitzen kann. Daß dieser Satz für unvollkommene Körper nicht allgemein richtig ist, wird an einem Beispiel gezeigt.

Um zur Konstruktion des Wurzelkörpers einer vorgegebenen Gleichung $f(x) = 0$ zu gelangen, wird nach einigen allgemeinen Sätzen über einfache, endliche, algebraische und transzendente Erweiterungen sowie über konjugierte Erweiterungen ein sogenannter Stammkörper konstruiert, das ist ein Körper, in dem die Gleichung $f(x) = 0$ eine Wurzel besitzt. Dieser Stammkörper erweist sich als Restklassenkörper, wie er im ersten Abschnitt betrachtet worden ist. Durch wiederholte Konstruktion von Stammkörpern kommt man zu einer kleinsten Erweiterung, wo $f(x)$ in Linearfaktoren zerfällt, und dieser Körper wird Wurzelkörper genannt und es wird schließlich die Eindeutigkeit des Wurzelkörpers nachgewiesen. Ein Paragraph über den sogenannten Fundamentalsatz der Algebra, daß im Körper aller komplexen Zahlen jeder Polynom mit reellen oder komplexen Koeffizienten in Linearfaktoren zerfällt, beschließt diesen wichtigen Abschnitt. Es wird darauf aufmerksam gemacht, daß dieser Satz einerseits nicht zur Algebra, sondern zur Funktionentheorie gehört, andererseits das Zentralproblem der Existenz des Wurzelkörpers zu einer gegebenen Gleichung dadurch nur für den Fall gelöst ist, daß die Koeffizienten der Gleichung Zahlen sind (rationale, reelle oder komplexe aus irgend einem Zahlkörper). Daß Steinitz zu jedem beliebigen abstrakten Grundkörper K die Existenz eines algebraisch abgeschlossenen Körpers A nachgewiesen hat von der Beschaffenheit, daß jedes Polynom mit Koeffizienten aus K in diesem Körper in Linearfaktoren zerfällt, wird bloß erwähnt. Für den Grundkörper der rationalen Zahlen ist dieses der Körper aller algebraischen Zahlen.

Im II. Abschnitt wird nun die Struktur der Wurzelkörper näher untersucht und vor allem nachgewiesen, daß für vollkommene Grundkörper mit unendlich vielen Elementen — die letztere Einschränkung ist nur für den Beweisgang, nicht aber für das Resultat erforderlich, wie Hasse auch betont — jeder Wurzelkörper (sowie überhaupt jede endliche algebraische Erweiterung) einfach ist. Der Wurzelkörper erweist sich aber weiter noch als Normalkörper oder Galoischer Körper. Die Automorphismen dieses Normalkörpers bilden die Galois'sche Gruppe der vorgelegten Gleichung und es wird schließlich

der Zusammenhang zwischen den Untergruppen der Galoisschen Gruppe und den Zwischenkörpern zwischen dem Grundkörper und dem Wurzelkörper erörtert.

Als Anwendung der Galoisschen Theorie wird im letzten Abschnitt schließlich die Frage nach der Auflösbarkeit algebraischer Gleichungen durch Wurzelzeichen behandelt, wobei das Resultat von Abel bewiesen wird, daß das allgemeine Polynom n -ten Grades über einen Körper K der Charakteristik Null für $n > 4$ nicht durch Wurzelzeichen auflösbar ist. Daß es auch im Körper der rationalen Zahlen Gleichungen gibt, deren Gruppe die symmetrische Gruppe ist, wird mit Heranziehung des Hilbertschen Irreduzibilitätssatzes bewiesen, während die Frage allgemein noch unentschieden ist.

Während in elementaren Lehrbüchern der Algebra fast durchwegs der Standpunkt eingenommen wird, daß als Koeffizientenkörper nur der Körper der rationalen oder reellen oder komplexen Zahlen zugelassen wird, stellt sich Hasse prinzipiell auf den Standpunkt der abstrakten Algebra als einer Theorie der abstrakt definierten Körper, Integritätsbereiche und Ringe und es ist wohl ganz außerordentlich zu begrüßen, daß dieser axiomatische Standpunkt einmal in einer Darstellung festgehalten wird, die sich an einen breiteren Leserkreis wendet. Vergeblich wird man daher in Hasses Darstellung alle jene Kapitel der Algebra suchen, die ganz wesentlich an die Voraussetzung gebunden sind, daß der Koeffizientenkörper der Körper der rationalen Zahlen ist. Ferner werden die Probleme außer acht gelassen, die sich auf mehrere Gleichungen höheren als ersten Grades mit mehreren Unbekannten beziehen. *T. Rella.*

Bieberbach-Bauer, Vorlesungen über Algebra. Unter Benutzung der dritten Auflage des gleichnamigen Werkes von G. Bauer in vierter, vermehrter Auflage, dargestellt von L. Bieberbach. Bei Teubner, Leipzig 1928, RM 20,—.

Wenn man Algebra als die Theorie rationaler Funktionen bezeichnet, so kann man dabei doch noch immer ganz verschiedene Dinge darunter verstehen. Handelt es sich um die Frage der Wurzelexistenz von Gleichungen mit Koeffizienten, welche reelle oder komplexe Zahlen sind, ebenso um Näherungsmethoden zur Bestimmung von Wurzeln von derartigen Gleichungen, so hat man es mit Fragen zu tun, welche kaum etwas mit den charakteristischen Methoden der Algebra zu tun haben und welche wohl eher ein Spezialgebiet der Funktionentheorie betreffen. Ganz anders sieht z. B. die Frage der Wurzelexistenz vom Standpunkt der abstrakten Körpertheorie aus und es möge bezüglich dieser prinzipiellen Stellungnahme auf die Werke über Algebra von Perron und Hasse hingewiesen werden.

In den Vorlesungen von Bieberbach-Bauer über Algebra hat man es nun fast ausschließlich mit Fragen zu tun, die dem ersten eben skizzierten Typus angehören und die ganz besonders schön und ausführlich behandelt werden und auch weitgehend Rücksicht auf die Bedürfnisse des numerischen Rechnens nehmen. Ferner möge auf die ausgezeichnete Darstellung der Determinantentheorie und der quadratischen und bilinearen Formen hingewiesen werden, die z. B. jedermann wärmstens zu empfehlen sind, der sich dem Studium der linearen Integralgleichungen zuwenden will.

Etwas stiefmütterlich sind dagegen die Fragen der Körpertheorie und der Galoisschen Theorie der Gleichungen behandelt. Für die allgemeine Gleichung n -ten Grades wird, wie in allen älteren Darstellungen, der Satz von der Wurzelexistenz aus dem Fundamentalsatz der Algebra entnommen, obwohl dieser doch wesentlich Zahlenkoeffizienten voraussetzt und nicht Unbestimmte.

T. Rella.

H. Beck, Einführung in die Axiomatik der Algebra. Bei de Gruyter, Berlin 1926.

Der Verfasser stellt sich zwei Aufgaben: erstens und hauptsächlich den Anfänger in das axiomatische Denken einzuführen, zweitens aber die Algebra der Matrizen genauer zu entwickeln. Zur Durchführung des ersten Programmpunktes geht Beck von der Gesamtheit der reellen Zahlen aus und den zwischen ihnen bestehenden Rechengesetzen und Anordnungssätzen, abstrahiert dann von konkretem Beiwerk und behält nur ein System von Dingen, zwischen denen die Relationsbeziehung der Äquivalenz und des Größer und Kleiner

definiert sind, ferner zwei Rechenoperationen und ihre inversen, die bestimmten formalen Gesetzen genügen. Für die Tätigkeit des Abstrahierens prägt er einen neuen Ausdruck aus dem vulgären Sprachgebrauch und nennt sie „pfeifen“. Zum erstenmal erscheint dieser Ausdruck, wegen dessen der Verfasser in einer Anmerkung um Nachsicht bittet, bei der Definition des Äquivalenzbegriffes. Es heißt dort wörtlich: „Für uns ist die Gleichheit nur synonym mit der gegenseitigen Vertretbarkeit; wir pfeifen auf alles übrige, aber auch auf alles. Diese Tätigkeit des Pfeifens ist sehr richtig; damit ist man bereits in die Traum- und Zaubersphäre der axiomatischen Walpurgisnacht eingegangen.“

Das zweite Kapitel zeigt am Beispiel von ebenen Punktmengen einen Bereich von Dingen auf, für den auch zwei Verknüpfungsoperationen definiert werden können, die teilweise die bei den Zahlen erkannten Gesetze befolgen, aber nicht alle und für den besonders die Anordnungssätze nur zum geringsten Teil aufrecht erhalten werden können. Beck nimmt diese Untersuchung vor allem aus dem Grund auf, um die Unabhängigkeit gewisser Axiome von den anderen zu erweisen.

Im dritten Kapitel werden Zahlenpaare behandelt. Nach Fallenlassen der Anordnungsaxiome entsteht einerseits bei einer Erklärung der Multiplikation als Hamiltonsche Paare (gemeine komplexe Zahlen) ein System von Dingen mit zwei Verknüpfungsoperationen, welches allen für reelle Zahlen bestehenden Gesetzen genügt, mit Ausnahme der Anordnungssätze, wodurch die Unabhängigkeit der Anordnungsaxiome von den Verknüpfungsaxiomen erwiesen ist. Andererseits ergeben sich bei anderer Definition der Multiplikation die Studyschen Paare, für welche das von Beck als „Gleichungssatz“ bezeichnete Axiom nicht gilt, daß ein Produkt nur dann verschwindet, wenn mindestens ein Faktor verschwindet, und damit ist die Unabhängigkeit dieses Axioms von den Verknüpfungsaxiomen nachgewiesen. Schließlich wird noch die von der gewöhnlichen, durchaus verschiedene Division im Bereich der Studyschen Paare untersucht.

Mit der Theorie der Matrizen im vierten Kapitel beginnt Beck die Frage nach der Unabhängigkeit des Axioms der Kommutativität der Multiplikation von den übrigen Verknüpfungsaxiomen zu untersuchen oder, wie er sagt, „die axiomatische Axt an die Wurzel noch eines anderen Grundsteines unseres Gebäudes zu legen, an das Kommutationsgesetz der Multiplikation. Daß $3 \cdot 4 = 4 \cdot 3$ ist, ist den Menschen aller Zeiten, ja, nach glaubwürdigen Berichten, sogar den Elberfelder Pferden klar gewesen. Und doch liegt hier kein logischer Zwang im vielfach erklärten Sinn vor, sondern wieder eine Materialeigenschaft. Der Beweis wird erbracht sein, wenn es gelingt, Objekte unserer Verknüpfungen aufzuzeigen, die den übrigen Verknüpfungsaxiomen gehorchen, nicht aber dem Axiom III 2.“ Dieses Heranziehen eines Beispiels aus der Lehre von den natürlichen Zahlen mit der Bemerkung „und doch liegt hier kein logischer Zwang vor“ erscheint dem Referenten mindestens bedenklich, da ja gerade für die natürlichen Zahlen ein Axiomensystem vorliegt, das der Verfasser im letzten Kapitel auch behandelt und aus dem nach Erklärung der Multiplikation die Kommutativität gefolgert werden kann, also wohl als logischer Zwang bezeichnet werden muß. Nachdem Beck festgestellt hat, daß bei üblicher Erklärung von Addition und Multiplikation von quadratischen Matrizen das Kommutationsgesetz der Multiplikation im allgemeinen nicht gilt, wohl aber für skalare Matrizen, beschließt er den betreffenden Paragraphen mit folgender, für ein mathematisches Lehrbuch wohl eigentümlich anmutender Redewendung: „Im Falle der Zahlen bedeutet III 2 nicht eine logische Notwendigkeit; es liegt hier eine Verstümmelung eines früheren stolzen Baues vor. Damals konnten Hero AB und Leander BA nicht zusammenkommen, das Wasser war viel zu tief. Jetzt sind aber die feindlichen Wasser außerhalb der Hauptdiagonale abkanalisiert, und in den Armen liegen sich beide!“ Diese Worte charakterisieren auch einigermaßen den Stil des Werkes. Sonderbar mutet wohl auch der Gedanke an, das System der reellen Zahlen eine Verstümmelung des Systems der Matrizen zu nennen. Die Frage nach der Möglichkeit der Division von Matrizen führt nun naturgemäß zur Theorie der linearen Gleichungen und

zu genauerem Studium von verschiedenen Äquivalenzbegriffen für Matrizen und diesen Dingen sind die folgenden Kapitel gewidmet.

Kapitel V handelt von Vektoren sowohl im dreidimensionalen wie im n -dimensionalen Raum und von Quaternionen, Kapitel VI von linearen Gleichungen, Kapitel VII von linearen Vektorgebilden, Kapitel VIII von bilinearen und quadratischen Formen, Kapitel IX vom Übergang zu homogenen Koordinaten, Kollineationen, Korrelationen und Hyperflächen zweiter Ordnung, Kapitel X schließlich von Determinanten und Tensoren. Bei der Herstellung der Normalform einer symmetrischen Matrix mittels kogredienter Transformationen liegt auf Seite 102 und 103 ein grobes Versehen vor, indem behauptet wird, daß

$$\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 0 & b_{12} \\ b_{21} & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} = \begin{pmatrix} b_{21} & 0 \\ 0 & b_{12} \end{pmatrix} \text{ ist.}$$

Mit dem Kapitel XI wird die Axiomatik fortgesetzt und zuerst über Gruppenaxiome gesprochen. Beck bedient sich folgender 4 Axiome zur Festlegung einer Gruppe:

G 1: Zu zwei Elementen A und B des Systems gibt es stets ein einziges Element A.B des Systems.

G 2: Die Komposition ist assoziativ (A.B).C=A.(B.C).

G 3: Aus A.C=B.C und ebenso C.A=C.B folgt stets A=B.

G 4: Zu zwei Elementen A und B des Systems gibt es stets ein einziges X und ein einziges Y im System so, daß B.X=A und X.B=A ist.

Beck behauptet, daß dies das Axiomensystem ist, wie es in Hecke, „Theorie der algebraischen Zahlen“ steht, doch ist der zweimalige Zusatz „einziges“ in G 4 sein Werk. Bei Hecke fehlt er richtigerweise, denn die Eindeutigkeit des Quotienten folgt nach G 3. Wird aber G 4 wie bei Beck formuliert, dann ist G 3, wie man fast unmittelbar sieht, eine Folge von G 4. Beck behauptet dagegen auf Seite 171 oben: „Aber auch umgekehrt ist deswegen (weil G 4 unabhängig von G 1, G 2 und G 3) auch das System der drei Axiome G 1, G 2, G 3 von G 4 unabhängig“ und verkündet nach einer sehr anfechtbaren Überlegung: „Damit ist in der Tat die Unabhängigkeit der vier Gruppenaxiome bewiesen.“ Er fährt fort: „Man könnte vielleicht einwenden, daß G 1 möglicherweise von G 2, G 3, G 4 abhängen könnte. Das würde heißen: aus der Gültigkeit von G 2, G 3, G 4 folgt die Gültigkeit von G 1; aus der Gültigkeit von G 2 und G 3 und der Nichtgültigkeit von G 4 müßte dann die Nichtgültigkeit von G 1 hervorgehen. Wir haben aber bereits in § 115 gesehen, daß das nicht der Fall ist.“ Ben Beweis für die zweite Behauptung bleibt Beck natürlich schuldig, denn sie ist falsch. Im Axiomensystem von Beck folgt z. B. aus G 1, G 2 und G 4 das Axiom G 3, dagegen kann aus der Gültigkeit von G 1 und G 2 und der Nichtgültigkeit von G 4 nicht auf die Nichtgültigkeit von G 3 geschlossen werden. Dieser § 119 über die Unabhängigkeit der Gruppenaxiome ist nach diesen Proben wohl als besonders mißglückt zu bezeichnen und dabei liegt doch gerade für diese Fragestellung z. B. die ausgezeichnete Darstellung von Loewy „Lehrbuch der Algebra“ vor, die Beck sogar in einer Fußnote dieses Paragraphen zitiert! Es folgt dann der Begriff des Körpers, der durch 10 Postulate in bekannter Weise festgelegt werden kann. In § 123 wird dann die Unabhängigkeit der 10 Axiome untersucht, aber in sehr unvollständiger Weise, indem jeweils nun die Unabhängigkeit eines Axioms von den vorausgegangenen konstatiert wird, wenn die Axiome in bestimmter Reihenfolge angenommen werden.

Den Abschluß des Buches bildet ein Kapitel über den genetischen Aufbau der Algebra, ausgehend von den Peanoschen Axiomen für natürliche Zahlen respektive ganze Zahlen, wobei für den versuchten Unabhängigkeitsbeweis der Peanoschen Axiome dieselben Bemerkungen gelten wie für den versuchten Unabhängigkeitsbeweis der 10 Körperpostulate.

T. Rella.

A. Speiser, *Theorie der Gruppen von endlicher Ordnung*. Zweite Auflage (Sammlung Courant). Bei Springer, Berlin 1927. Geb. RM 16,50.

Diese Auflage weist gegenüber der ersten große, durchwegs außerordentlich begrüßenswerte Veränderungen auf. Da sind einmal zwei kleine Aufsätze

in die Einleitung aufgenommen, die einerseits gruppentheoretische Fragen als einen Hauptgegenstand antiker Mathematik in der Theorie der regelmäßigen Figuren aufweisen, andererseits den Gruppenbegriff ableiten aus dem Begriff des Gruppoids; dieser Name stammt wohl erst von Brandt, die darin zum Ausdruck kommende mathematische Abstraktion wird nachgewiesen beim Aufstieg von der Streckenaddition zur Vektoraddition, von der Ordinalzahl zum Differenzsymbol (Speiser sagt zur Kardinalzahl, was dem Referenten bedenklich vorkommt), vom Übergang eines Anfangs- in einen Endzustand zum Begriff der Zeit, vom Zweiklang zum Intervall und wird schließlich ausführlich dargestellt am Begriff der Permutationen.

Ausführlicher dargestellt sind die Abschnitte über Abelsche und Substitutionsgruppen. Ganz besondere Erweiterungen hat aber das Kapitel über Kristallstrukturen erfahren, indem auch die entsprechenden ebenen Probleme behandelt werden, auf deren Bedeutung für die antike Mathematik bereits in der Einleitung hingewiesen ist. Das Kapitel über Symmetrien der Ornamente, das auch als Einführung in den Gruppenbegriff vorgenommen werden kann, erscheint dem Referenten besonders geglückt und sehr dankenswert die Aufnahme vieler Illustrationen aus dem Werk von Owen Jones „Grammatik der Ornamente“.

Neben diesen bedeutenden Erweiterungen ist der Text sorgfältig revidiert und verschiedene kleine Mängel der ersten Auflage erscheinen beseitigt.
T. Rella.

L. E. Dickson, Algebren und ihre Zahlentheorie mit einem Kapitel über Idealtheorie von A. Speiser. Deutsche Übersetzung von J. J. Burckhardt und E. Schubarth. (Veröffentl. d. Züricher Math.-Ges., Bd. 4.) Bei O. Füssli, Zürich 1927. Geh. Fr. 18,—, geb. Fr. 22,— (RM 17,60).

Es ist außerordentlich zu begrüßen, daß durch diese Übersetzung Dicksons grundlegendes Werk einem größeren deutschen Publikum leichter zugänglich gemacht worden ist. Die neueren Theorien der hyperkomplexen Zahlen (wie dieser Problemkreis gewöhnlich in Deutschland bezeichnet wird) unterscheiden sich grundlegend von den älteren Theorien dadurch, daß als Grundbereich ein beliebiger Körper zugelassen wird und nicht die Einschränkung auf den Körper aller reellen oder aller komplexen Zahlen gemacht wird. Dadurch hat die Theorie ganz wesentlich an Allgemeinheit gewonnen und umfaßt die alte vollständig. Sehr unangenehm wird wohl von deutschen Lesern empfunden werden, daß nur die aus dem Englischen übernommenen Kunstausdrücke, nicht aber auch die im Deutschen gebräuchlichen Fachausdrücke erwähnt werden. Das Buch erfordert zwar anscheinend keine besonderen Vorkenntnisse, trotzdem wird es kaum mit Vorteil von jungen Leuten gelesen werden können, die nicht schon in der Gruppentheorie und Galoisschen Theorie einige Vorkenntnisse besitzen.

Das I. Kapitel bringt nur ein einleitendes Referat über Körper, Polynome und Matrizen, Kapitel II die Einführung in die Algebren, wobei in der ursprünglichen Definition die Assoziativität der Multiplikation nicht gefordert wird. Allerdings sind nur die ersten Kapitel auch für nicht assoziative Algebren von Bedeutung, während vom VI. Kapitel an die Assoziativität der Multiplikation wesentlich gefordert wird. Schon in diesem Kapitel wird die Äquivalenz jeder assoziativen Algebra mit einer Matrixalgebra nachgewiesen. Kapitel III behandelt den Spezialfall der Divisionsalgebren, das sind assoziative Algebren ohne Nullteiler, wobei der bekannte Satz von Frobenius bewiesen wird, daß die einzigen Divisionsalgebren im Körper der reellen Zahlen dieser Körper selbst, der Körper aller komplexen Zahlen und die Algebra der reellen Quaternionen sind. Für genügend allgemeine Grundkörper ergeben sich dagegen unendlich viele Divisionsalgebren. Kapitel IV handelt von Linearsystemen einer Algebra (Modeln), Kapitel V entwickelt die fundamentalen Begriffe invariante Subalgebra, direkte Summe, Reduzibilität und Differenzenalgebra. Mit Kapitel VI wird die Untersuchung über die Struktur der Algebra durch Untersuchung der nilpotenten und idempotenten Größen einer Algebra begonnen (nilpotent heißt eine GröÙe, wenn eine genügend hohe

Potenz verschwindet, idempotent eine Größe e , für die $e^2 = e$ gilt). Diese Untersuchungen werden in Kapitel VII zu Ende geführt und ergeben den Hauptsatz, daß jede einfache Algebra (Algebra ohne invariante eigentliche Subalgebra) das direkte Produkt einer Divisionsalgebra und einer vollständigen Matrixalgebra ist, und umgekehrt und weiters den Satz, daß jede assoziative Algebra in einem Körper K entweder halbeinfach ist (Algebra ohne invariante eigentliche nilpotente Subalgebra) oder die Summe ihres Radikals und einer halbeinfachen Subalgebra, deren jede in K liegt.

Mit dem Kapitel VIII beginnt der zahlentheoretische Teil, und zwar bringt dieses Kapitel eine kurze Darstellung einiger bekannter Tatsachen aus der Theorie der algebraischen Zahlen. Kapitel IX ist der Zahlentheorie der verallgemeinerten Quaternionenalgebren gewidmet samt Anwendungen auf die Theorie der rationalen Zahlen. Dieses Kapitel ist besonders ausführlich gehalten (44 Seiten). Kapitel X bringt die allgemeine Zahlentheorie der Algebren, wobei wesentlich die Definition der ganzen Größen einer Algebra A ist, als Elemente eines willkürlichen, aber fest gewählten größten Integritätsbereiches J von Größen aus A , der die Haupteinheit enthält und welche zudem ganzen Gleichungen genügen. Dabei ist unter Integritätsbereich ein Bereich J von Größen einer Algebra A verstanden, der so beschaffen ist, daß er mit zwei Größen auch ihre Summe, Differenz und Produkt enthält, ohne daß deforziert wird, daß ein Produkt nur verschwindet, wenn mindestens ein Faktor verschwindet. Kapitel XI entwickelt kurz die Definition und einige Sätze aus der Theorie der abstrakten Körper, Quasikörper (diese unterscheiden sich von Körpern dadurch, daß die Multiplikation nicht als kommutativ vorausgesetzt ist, dagegen soll jedes Element ein inverses besitzen) und speziell der endlichen Körper (Galoisfelder) und es wird der Satz bewiesen, daß jeder endliche Quasikörper ein Körper ist. Kapitel XII gibt in Ergänzung der Kapitel I—VII solche Körper an, in denen die dort entwickelten Sätze auch gelten. Kapitel XIII schließlich bringt als Anhang die Idealtheorie in rationalen Algebren, die von Speiser im 71. Band der Vierteljahrsschrift der Naturforschenden Gesellschaft in Zürich 1926 publiziert worden ist. *T. Rella.*

H. Graßmann, Projektive Geometrie der Ebene, unter Benutzung der Punktrechnung dargestellt. Zweiter Band: Ternäres. Zweiter Teil. 522 Seiten. Bei Teubner, Leipzig 1927. Geh. RM 20,—, geb. RM 22,—.

Nach Behandlung der projektiven Geometrie auf einer Kurve zweiter Ordnung und zweiter Klasse wird das Scharbüschel von Kurven zweiter Ordnung und das hierzu duale Gebilde untersucht. Außerordentlich eingehend wird die Reziprozität studiert, nachdem die Kollineationen ermittelt wurden, die eine Büschelschar und ein Scharbüschel in sich überführen. In den Abschnitten über die Reziprozität wird man mit Interesse die Behandlung der Probleme mittels der extensiven Brüche lesen.

Bemerkenswert ist die ausführliche Untersuchung der Apolarität. An dieser Stelle finden die Lückenformen zweiter Ordnung weitgehende Verwendung.

Die projektive Geometrie der Kurven dritter Ordnung und der Kurven dritter Klasse und damit auch die Untersuchung der entsprechenden Lückenausdrücke ist sehr breit angelegt und bietet reichlich Gelegenheit, den Grassmannschen Kalkül zu zeigen.

Die Ergebnisse dieser Abschnitte finden bei der Entwicklung der Geometrie des Kegelschnittnetzes und des Kegelschnittgewebes sowie der Zusammenhänge zwischen diesen Gebilden und den ihnen zugeordneten Hesseschen und Cayleyschen Kurven weitgehende Verwendung.

In den bisher erwähnten Abschnitten fanden metrische Fragen nur ganz gelegentlich Erwähnung. Systematisch wird die Metrik in den Schlußabschnitten entwickelt. Hier sind besonders zu erwähnen: die Einführung des Polarsystemes, des Kreispunktpaares und die systematische Erledigung der einschlägigen Fragestellungen mittels der Punktrechnung. Bemerkenswert scheint auch die Besprechung der projektiv oder metrisch ausgezeichneten Kegelschnittnetze und -gewebe.

Eine eingehendere Aufzählung aller in diesem so außerordentlich klaren Buche entwickelten Probleme ist nicht möglich (das Inhaltsverzeichnis allein um-

faßt 10 Seiten!). Der Leser, der das Buch zum Nachschlagen heranzieht, wird die breite Darstellung angenehm empfinden. Die Geometer, die sich mit projektiven Fragen beschäftigen, werden dem Verlage und dem Herausgeber sicher für diese umfassende Publikation danken. *Knoll.*

R. Baldus, Nichteuklidische Geometrie. (Hyperbolische Geometrie der Ebene.) Sammlung Götschen, Nr. 970. Bei de Gruyter, Berlin 1927. 152 Seiten. RM 1,50.

Eine sehr klare und verständliche Einführung in die Grundtatsachen der hyperbolischen Geometrie. Nach einer historischen Einleitung wird zur Grundlegung der absoluten Geometrie das System der Hilbertschen Axiome eingeführt, wobei das Parallelenaxiom abgesondert nach den übrigen Axiomen besprochen wird. Die Möglichkeit der hyperbolischen Geometrie wird nun an der auf einen reellen Kreis gegründeten projektiven Maßbestimmung gezeigt. Und dasselbe Modell dient nun zur Herleitung des ganzen Komplexes der Sätze der hyperbolischen Geometrie. Die Beweisführungen und Ableitungen sind im einzelnen sehr schön durchgeführt, wobei bald rein projektiv gearbeitet, bald die Einbettung der projektiven Maßbestimmung in die Euklidische Ebene benützt wird. Den Schluß des Buches bilden einige Ausblicke auf die Möglichkeit anderer Maßbestimmungen und die Beziehungen zur Wirklichkeit. Das Buch ist dem Zweck der Sammlung Götschen entsprechend, ohne große Vorkenntnisse vorauszusetzen, sehr verständlich, sowie auch methodisch sehr einheitlich. Einen Einwand, der jedoch den mathematischen Wert des Buches kaum beeinträchtigt, möchte der Referent nicht unterdrücken. Durch die gewählte, auf Felix Klein zurückgehende Form der Darstellung, in welcher die Sätze der n. E.-Geometrie konsequent als Sätze der auf einen Kegelschnitt gegründeten Maßbestimmung abgeleitet werden, ergibt sich zwar eine äußerst konzise Beweisführung, allein es scheint, daß auf diese Weise der ursprüngliche geometrische Gehalt der n. E.-Geometrie als Analyse eines möglichen Erfahrungsraumes speziell für den Anfänger etwas verdunkelt wird. *Helly.*

Hefter-Koehler, Lehrbuch der analytischen Geometrie. Bd. I, zweite Auflage. Bei Braun, Karlsruhe 1927. Geh. RM 20,—.

Der wichtigste Unterschied der zweiten Auflage gegenüber der ersten besteht darin, daß nun dem Buch ein einleitendes Kapitel „Die Grundlagen der Geometrie“ vorangeschickt wurde, welches die Axiomatik der projektiven und der euklidischen Geometrie behandelt. Damit ist, wie es in der Vorrede heißt, „das ganze Werk axiomatisch fundiert und alles ausgemerzt, was früher noch der Elementargeometrie und der Anschauung entnommen werden mußte“. Das letztere Bestreben ist natürlich sehr zu begrüßen. Aber speziell einer axiomatischen Fundierung bedarf, nach Ansicht des Referenten, die analytische Geometrie ebensowenig, wie die Axiomatik einer analytischen Fundierung. Wenn man erstens Axiomatik der Elementargeometrie (d. h. der Lehre von den einfachen Raumgebilden, wie Punkte, Geraden, Ebenen usw.) betreibt, so geht man von einem System von Dingen aus, über deren Natur nichts vorausgesetzt wird und die nur der suggestiven Ausdrucksweise halber als Punkte, Gerade, Ebenen usw. bezeichnet werden, während über das ganze System vorausgesetzt wird, daß es gewissen Bedingungen (den Axiomen) genügt. Man untersucht sodann, was man über das System aus den vorausgesetzten Axiomen schließen kann. Insbesondere führt man die Untersuchung bis zu einem Punkt, auf welchem man zeigen kann: Es ist möglich, das betrachtete System der Punkte auf die Menge aller geordneten n -Tupel von reellen Zahlen eindeutig und derart abzubilden, daß den Geraden, Ebenen usw. Mengen von n -Tupeln, welche linearen Gleichungen genügen, entsprechen. Wenn man zweitens analytische Geometrie betreibt, dann geht man, wie man seit Study weiß, von der Menge der geordneten n -Tupel reeller Zahlen aus, bezeichnet die Menge der n -Tupel, welche linearen Gleichungen genügen, als Geraden, Ebenen usw. und untersucht, was sich über diese Gebilde und ihre gegenseitigen Beziehungen errechnen läßt. Insbesondere führt man die Untersuchung so weit, daß man einsieht, daß die bei der axiomatischen Behandlung der Geometrie postulierten Beziehungen er-

füllt sind. Das einzige Fundament der analytischen Geometrie ist also die Lehre von den reellen Zahlen, die sich übrigens axiomatisch (mit Hilfe von Anordnungsaxiomen, also von einem Teil der elementar-geometrischen Axiome) begründen läßt. Was aber das Verhältnis der analytischen und der axiomatischen Geometrie betrifft, so stellen die beiden zwei methodisch ganz verschiedene Behandlungsweisen der Geometrie dar, deren gegenseitige Beziehungen in den Details zu mathematisch sehr interessanten Problemen und Resultaten führt, — man denke an die Untersuchungen über die Sätze von Desargues und Pascal —, wobei es aber keinen Sinn hat, zu sagen, daß eine der beiden Behandlungsweisen eine Fundierung der anderen liefere. Die Durchführung des Studyschen Standpunktes in der neuen Auflage wäre erheblich vorteilhafter gewesen als die tatsächlich vorgenommenen Änderungen. *Menger.*

H. Beck, Elementargeometrie. Akademische Verlagsgesellschaft, Leipzig 1929. —

Das Buch bringt eine Einführung in die analytische Geometrie der Ebene und beschränkt sich auf die Geometrie der Anschauungsebene, vermeidet also alle unendlich fernen und imaginären Elemente. Es beginnt mit einfachen Abbildungen, insbesondere offenen Abbildungen. Hierauf werden Abbildungsgruppen im allgemeinen und die Gruppen der Dehnungen und Bewegungen im besonderen besprochen. Ausführlich wird die analytische Affingeometrie behandelt. Dann folgt ein kurzer Abschnitt über Statik. Im letzten Kapitel wird der Winkel und die Entfernung eingeführt, ohne daß dabei die unendlich fernen Elemente zu Hilfe genommen werden.

Ein klarer, übersichtlicher Aufbau macht das Buch leicht verständlich. Es kann besonders Hochschülern in den ersten Semestern, aber auch Mittelschullehrern sehr empfohlen werden. *Hofreiter.*

L. P. Eisenhardt, Non-Riemannian Geometrie. (New York, American Mathematical Society, 1927.) Dollar 2,50.

Die Christoffelklammern $\left\{ \begin{smallmatrix} p \\ r \end{smallmatrix} q \right\}$, die keine Tensoren sind, befolgen Transformationsgesetze, die wieder ihrerseits verwendet werden können zur Definition von Größen Γ_{pq}^r , die also (wie die $\left\{ \begin{smallmatrix} p \\ r \end{smallmatrix} q \right\}$) von drei Indizes abhängen, und zu denen die $\left\{ \begin{smallmatrix} p \\ r \end{smallmatrix} q \right\}$ des Riemannschen Raumes gehören. Für die Γ_{pq}^r , die in den unteren Indizes nicht symmetrisch zu sein brauchen, ist charakteristisch, daß

$$d\lambda^i + \Gamma_{pq}^i \lambda^p dx^q$$

mit λ^i ein kontravarianter Vektor ist.

Will man die Eigenschaften eines Riemannschen Raumes untersuchen, die nicht rein metrischer Natur sind, so adjungiert man den betrachteten Punktraum eine „Drei-Indizes-Größe“ Γ_{qp}^r , und stellt die Simultaninvarianten

der geometrischen Größen und des Γ_{pq}^r her. Die so entstehende „Nicht-

Riemannsche“ Geometrie könnte man positiver wohl als die Geometrie der „Parallelverschiebung“ bezeichnen.

Bei dem vorliegenden Buche möchte ich besonders die zweckmäßige Verwendung der Systeme totaler, nicht vollständig integrierbarer Differentialgleichungen hervorheben zur Erledigung einiger fundamentaler Probleme, wie der Existenz eindeutig parallel verschobener Vektoren, der Existenz und Eindeutigkeit eines Maßtensors zu vorgegebener Parallelverschiebung usw. Zu den Vorteilen des Buches zählt noch die ungemein klare Schreibart des Autors, so daß auch dem deutschen Leser keine Schwierigkeiten des Verhältnisses entstehen.

W. Mayer.

Enzyklopädie der mathematischen Wissenschaften. Redigiert von W. Fr. Meyer und H. Mohrmann. Bei Teubner, Leipzig 1902/27. RM 44,20.

3. Teil des dritten Bandes (Gnometrie) ist nunmehr abgeschlossen. Die Enzyklopädie bedarf wohl keiner weiteren Empfehlung, eine Inhaltsangabe des Bandes genügt *):

Heft 1 (14. X. 1902): Mangold: Anwendung der Differential- und Integralrechnung auf Kurven und Flächen. Lilienthal. Die auf einer Fläche gezogenen Kurven.

Heft 2/3 (20. IX. 1903): Scheffers: Besondere transzendente Kurven. Lilienthal: Besondere Flächen. Voss: Abbildung und Abwicklung zweier Flächen aufeinander.

Heft 4 (14. V. 1915): Liebmann: Berührungstransformationen. Derselbe: Geometrische Theorie der Differentialgleichungen.

Heft 5 (8. II. 1921): Salkowski: Dreifach orthogonale Flächensysteme.

Heft 6 (1. XI. 1922): Weitzenböck: Neuere Arbeiten der algebraischen Invariantentheorie. Differentialinvarianten.

Heft 7 (15. X. 1927): Berwald: Differentialinvarianten in der Geometrie. Riemannsche Mannigfaltigkeiten und ihre Verallgemeinerungen. Inhaltsverzeichnis und Register. *W. Mayer.*

F. Baur, Korrelationsrechnung. (Math.-phys. Bibliothek 75). Bei Teubner, Leipzig u. Berlin 1928. RM 1,20.

Die Lehre von der Korrelation, diese jüngste Weiterbildung der Wahrscheinlichkeitsrechnung, ist dazu berufen, die Anwendung der Mathematik auch auf Gebiete, die ihr bisher so ziemlich verschlossen waren, nämlich auf die beschreibenden Naturwissenschaften, zu ermöglichen. Begreiflicherweise ergibt sich daraus der Wunsch, den Zugang zu diesem Zweig zu eröffnen. Diesem Zwecke dient das kleine Buch von Baur. Es schließt sich in der Darstellung ziemlich eng an das vor drei Jahren erschienene Buch von Tschuprow an und kann insbesondere auch mit Vorteil als Vorbereitung zu diesem Verwendung finden. *Schrutka.*

A. Haas, Materiewellen und Quantenmechanik. Eine Einführung auf Grund der Theorien von De Broglie, Schrödinger, Heisenberg und Dirac. Zweite, verbesserte und wesentlich vermehrte Auflage. VI + 179 Seiten. Mit 3 Abbildungen. Bei der Akad. Verlagsgesellschaft, Leipzig 1929. Geh. RM 7,—, geb. RM 8,—.

Da die erste Auflage des Büchleins in Band XXXV der Monatshefte, Literaturbericht, Seite 65/66, bereits ausführlich besprochen wurde, darf man sich darauf beschränken, auf die Änderungen der zweiten Auflage gegenüber der ersten einzugehen. Trotzdem zwischen dem Erscheinen der ersten und zweiten Auflage nur ein Zeitraum von einem halben Jahre liegt, mußte die zweite Auflage um vier neue Kapitel vermehrt werden, die den inzwischen erfolgten Fortschritten der Quantenmechanik gewidmet sind. Diese vier neuen Kapitel beinhalten 1. die Einwirkung von Licht- und Materiewellen auf Atome (darunter den erst kürzlich aufgefundenen Smekal-Raman-Effekt); 2. das Phänomen der quantenmechanischen Resonanz, das nach Heitler und London auch die Erklärung der homöopolaren Bindungen ermöglicht, welche der älteren Quantenmechanik unüberwindliche Schwierigkeiten bereitete; 3. die relativistische Verallgemeinerung der wellenmechanischen Grundgleichung als Vorbereitung zu 4. der Diracschen Theorie des Elektrons, die rein deduktiv das Phänomen des Elektronenspins liefert. Wiederum zeigt sich die Professor Haas eigene Darstellungskunst bei diesen neuesten und sehr schwierigen Errungenschaften der Theorie. *Sexl.*

A. Haas, Die Welt der Atome. Zehn gemeinverständliche Vorträge. Mit 37 Figuren im Text und auf 3 Tafeln. 8°. XII und 130 Seiten. Bei de Gruyter & Co., Berlin 1926. Geh. RM 4,80, geb. RM 6,—.

*) Das Datum in der Klammer bezieht sich auf die Ausgabe der einzelnen erschienenen Hefte dieses Bandes und nicht auf den Abschluß der einzelnen Artikel.

Das Büchlein besteht aus 10 Vorträgen, die für Hörer aller Fakultäten an der Wiener Universität gehalten wurden. Sie bezwecken, die Errungenschaften der modernen Atomphysik einem Laienpublikum in möglichst knapper und doch erschöpfender und zugleich leicht verständlicher Form darzulegen. Der Inhalt der zehn Vorträge ist im einzelnen: Materie und Elektrizität, Die Bausteine der Atome, Die Quanten des Lichtes, Spektren und Energiestufen, Das Wasserstoffatom, Die Grundstoffe, Das Atom als Planetensystem, Die Molekeln, Die Radioaktivität, Die Umwandlung der Grundstoffe. Am Schluß sind in ergänzenden Anmerkungen die markantesten Formeln und eine chronologische Übersicht beigelegt. Es erübrigt, dieses Buch aus der Feder von Artur Haas, das selbstverständlich auch die oft und vielgerühmten Vorzüge aller der anderen weitverbreiteten und bekannten Bücher des Autors besitzt, noch besonders empfehlen zu wollen. *Schl.*

H. Greinacher, Ionen und Elektronen. (Abhandlungen aus dem Gebiete der Mathematik, Naturwissenschaft und Technik, Heft 9.) Mit 26 Figuren. 8°. 58 Seiten. Bei Teubner, Leipzig 1924. Geh. RM 1,60.

Das Büchlein ist als Ergänzung zu des Verfassers „Einführung in die Ionen- und Elektronenlehre der Gase“, die aus Experimentalvorlesungen hervorgegangen ist, gedacht und behandelt in vorzugsweise theoretischer Form einzelne spezielle Kapitel eingehender. Es gliedert sich in 6 Abschnitte, die im einzelnen die Volumenionisierung, die Messung der Ionenströme, die Eigenschaften der Ionen, die Oberflächenionisierung, die Gesetze frei bewegter Ionen und die Verwendung der Elektronenröhre behandeln. *Schl.*

E. Lohr, Atomismus und Kontinuitätstheorie in der neuzeitlichen Physik. (Wissenschaftliche Grundfragen. Philosophische Abhandlungen, herausgegeben von R. Hönigswald, Band VI.) 8°. 82 Seiten. Bei Teubner, Leipzig 1926. Geh. RM 4,—.

In erweiterter Form erscheint diese Schrift als Wiedergabe einer Vortragsreihe gelegentlich der Salzburger Hochschulkurse im September 1925, die das Ziel hatte, Korpuskulartheorie und Kontinuitätstheorie in ihrer historischen Entwicklung und in ihrer grundsätzlichen Einstellung gegeneinander abzuwägen. Von der methodischen Überlegenheit der Kontinuitätstheorie, insbesondere in der modernen Form, die der berühmte Lehrer des Verfassers G. Jaumann geschaffen hat, überzeugt, bemüht sich der Verfasser gleichwohl redlich, auch der atomistischen Vorstellungs- und Gedankenwelt gerecht zu werden. Als Einführung und Darstellung des Jaumannschen Standpunktes, von dem bisher eine an breitere Kreise gerichtete Darstellung gefehlt hat, ist die vorliegende Schrift sehr zu begrüßen: Mancher Physiker wird zum ersten Male durch diese Schrift von der Jaumannschen Deutung der Kathoden- und Kanalstrahlen als Longitudinalwellen erfahren, durch die die wellentheoretische Auffassung der Materie vorweggenommen wurde, wenn auch dieser Gedanke Jaumanns den Entwicklungsgang der Physik nicht zu beeinflussen vermochte und von der de Broglieschen Konzeption der quantentheoretischen Phasenwellen grundverschieden ist. Es könnte scheinen, als ob durch diese Entwicklungen de Broglies auch das vom Verfasser befürchtete Nicht-Wieder-Zurückkehren-Können zum Kontinuitätsstandpunkte nun doch wieder ermöglicht und wünschenswert werde. Doch sind in den neuesten Gedankengängen der Quantentheoretiker die Korpuskular- und die Kontinuitätstheorie keine einander bekämpfenden Gegner mehr, sondern beide Standpunkte stellen gewissermaßen nur zwei gleichberechtigte Grenzgebiete dar, deren Synthese teilweise schon gelungen, aber im großen und ganzen noch eine überaus wichtige Aufgabe der im Gange befindlichen Forschung ist. So scheint sich der Gegensatz hie Korpuskular-, hie Kontinuitätstheorie in einer überlegenen, beide Standpunkte umfassenden Synthese aufzulösen, wie seinerzeit die Kantische Philosophie die einander bekämpfenden Richtungen des Realismus und des Idealismus in sich vereinigen wollte. *Schl.*

L. Schlesinger, Raum, Zeit und Relativitätstheorie. (Abhandlungen aus dem Gebiete der Mathematik, Naturwissenschaft und Technik, Heft 5.) Mit 7 Figuren, 40 Seiten. Bei Teubner, Leipzig 1920. Geh. RM 1,10.

Das Büchlein behandelt in allgemein verständlicher Form die Grundlehren der speziellen und allgemeinen Relativitätstheorie. Zur Erläuterung werden vorwiegend graphische Methoden benutzt. *Sezl.*

W. Müller, Dynamik. Band I. Dynamik des Einzelkörpers. (Sammlung Götschen, Band 902.) Band II. Dynamik von Körpersystemen. (Sammlung Götschen, Band 903.) Bei de Gruyter & Co., Berlin 1925. In Leinen je RM 1,25.

Die beiden Bändchen enthalten die Grundlehren der Dynamik in klarer, wenn auch knapper Form. Der erste Band gliedert sich in die Kinematik des Punktes, die Bewegung des Massenmittelpunktes und die vollständige Bewegung des starren Körpers. Der zweite Band ist unterteilt in das System und seine Kräfte, analytische Methoden der Systemdynamik und in prinzipielle Betrachtungen (= Variationsprinzipien, Hamilton-Jacobische Gleichung). Dem ersten Band ist ein ausführliches Literaturverzeichnis vorangestellt. *Sezl.*

F. R. Lipsius, Wahrheit und Irrtum in der Relativitätstheorie. 8^o. VIII und 154 Seiten. Bei J. C. B. Mohr, Tübingen 1927. Geh. M 7,50. Inhalt: Einleitung, Die spezielle Relativitätstheorie, Äther und Materie, Die Relativierung der Kausalität, Die allgemeine Relativitätstheorie, Literaturverzeichnis. *Sezl.*

Physikbüchlein. Ein Jahrbuch der Physik. Herausgegeben von Dr. Werner Bloch. Mit 45 Figuren. 80 Seiten. Bei Francksche Verlagshandlung, Stuttgart 1926. Geh. RM 1,50.

Inhalt: Der physikalische Ertrag des Jahres 1925 von Dr. Werner Bloch, Die Bestimmung des Elektrons von Dr. Kurt Gehlhoff, Die Theorie des Magnetismus von Dipl.-Ing. Dr. Robert Usmann, Die Geschichte des Thermometers von Fritz Löw, Die Bewegung der drei Körper von Prof. Dr. Kirchberger, Der Heizwert von Brennstoffen von Dr. Walter Block, Neuere Waagenkonstruktionen von Dr. Walter Block, Die Lufthülle der Erde von Dr. Alfred Booss, Optisches Glas von Dr. Georg Jäckel, Der Skin-Effekt von Dipl.-Ing. Dr. Robert Usmann. Alle Aufsätze in allgemeinverständlicher Form. Bemerkenswert sind die schönen, instruktiven Abbildungen zu dem Aufsatz über optisches Glas. *Sezl.*

H. Kafka, Die ebene Vektorrechnung und ihre Anwendung in der Wechselstromtechnik. Teil I: Grundlagen. (Sammlung mathematisch-physikalischer Lehrbücher, herausgegeben von E. Trefftz, Band 22.) Mit 62 Figuren. 8^o. VIII und 132 Seiten. Bei B. G. Teubner, Leipzig und Berlin 1926. Geh. M 7,60.

Der erste Teil dieses aus Hochschulvorlesungen entstandenen Werkes befaßt sich mit den Grundlagen der ebenen Vektorrechnung und ihren Anwendungen in der Wechselstromtechnik. Die Darstellung beschränkt sich auf stationäre Vorgänge und setzt zeitlich sinusförmig veränderliche Größen voraus. Die Grundlagen der ebenen Vektorrechnung werden möglichst ausführlich und anschaulich dargestellt, da die bisher verhältnismäßig geringe Verwendung des Wechselstromtechnik so gut angepaßten Werkzeuges wohl darauf zurückzuführen ist, daß ein Lehrbuch fehlt, das die Grundlagen der ebenen Vektorrechnung mit genügender Breite behandelt. Um die Anschaulichkeit der Schreibweise zu erhöhen, hat Verfasser für die Kennzeichnung von Drehungen Rundpfeile als Symbol eingeführt. Auch die Grundlagen der Maxwell'schen Theorie werden zur besseren Brauchbarkeit des Buches in einem Kapitel kurz entwickelt. Der Inhalt der sechs Kapitel ist im einzelnen folgender: Einführung in die ebene Vektorrechnung, Oszillatoren und Zeitvektoren, Zusammenstellung der für die Anwendungen der ebenen Vektorrechnung erforderlichen Grundlagen aus der Theorie der Elektrizität und des Magnetismus, Selbstinduktion und Kapazität im Wechselstromkreis, Parallelschaltung von Wechselstromzweigen, Ström- und Leistungskomponenten. Das Buch wird bei Elektrotechnikern sicherlich die große Verbreitung finden, die ihm infolge seiner Anschaulichkeit und Brauchbarkeit gebührt. *Sezl.*

P. Bommersheim, Beiträge zur Lehre von Ding und Gesetz. Wissenschaftliche Grundfragen, herausgegeben von R. Hönigswald, Heft 8. Bei Teubner, Leipzig 1927. VIII u. 108 S.

Die Hauptaufgabe dieser Untersuchungen, die erkenntnistheoretisch auf neukantischem Standpunkt stehen, ist eine logische: für die Begriffe Ding, Eigenschaft, Kausalgesetz, Naturgesetz sollen Definitionen aufgestellt und Einteilungen in Unterarten gegeben werden; Hauptfrage: „Zu welcher Struktur werden die Naturgesetze im Begriff des Dinges verbunden?“. Die Lösung der Aufgabe wird mit Sorgfalt und Scharfsinn versucht; trotzdem finden sich in den Ergebnissen viele Unzulänglichkeiten und Schiefheiten, hauptsächlich infolge der völligen Vernachlässigung der in den letzten Jahrzehnten entwickelten neuen Logik. *Carnap.*

Riebesell, Die Relativitätstheorie im Unterricht. 5. Beiheft der Unterrichtsblätter f. Math. u. Naturw. Bei Salle, Berlin 1926.

In dieser Schrift wird gezeigt, welcher Stoff der Relativitätstheorie im mathematischen und physikalischen Unterricht an höheren Lehranstalten behandelt werden kann. Sehr ausführlich wird die spezielle Relativitätstheorie im Falle einer Relativbewegung zweier Systeme parallel zu einer Achse erörtert. Als besonders gelungen können das 6. und 7. Kapitel bezeichnet werden, in welchen die Folgerungen aus den Transformationsgleichungen und die Veranschaulichung der angegebenen Resultate angegeben werden.

J. Peters, Die Grundlagen der Musik. VIII + 156 Seiten. Bei Teubner, Leipzig und Berlin 1927. RM 7,60.

Seinem in der math.-physik. Bibl. als Bd. 55 erschienenen Büchlein, betitelt: „Die mathematischen und physikalischen Grundlagen der Musik“ ließ der Verfasser das vorliegende Werk folgen. Es behandelt denselben Gegenstand in bedeutend breiterer Darstellung, außerdem aber noch die physiologischen und psychologischen Grundlagen der Musik. An Vorkenntnissen wird nie mehr vorausgesetzt, als was der allgemeinen, in der Mittelschule erworbenen Bildung entspricht. Die neuesten Fragestellungen werden angeschnitten. Das Buch ist äußerst klar geschrieben und kann jedem Musiker aufs beste empfohlen werden.

J. Lense.

Meyer und Braun. Arithmetik und Algebra. Arbeitsbuch für die Mittelklassen der Oberlyzeen und Studienanstalten, die Lyzeen und höheren Mädchenschulen. Herausgegeben von J. Inkmann, C. Müller, W. Schwarz und J. Witte. Aschendorffsche Verlagsbuchhandlung, Münster i. W., 1928.

Das Lehrbuch umfaßt die Operationen 1., 2. und 3. Stufe im Gebiete der rationalen und der reellen Zahlen, das Rechnen mit komplexen Zahlen, die Gleichungen 1. und 2. Grades, sowie die einfachsten algebraischen Funktionen. Die Rechengesetze werden an der Hand besonderer Zahlenbeispiele gewonnen; der Forderung der Verwendung des Funktionsbegriffes wird durch eine kurze Behandlung der linearen und quadratischen Funktion und deren Anwendung auf die graphische Auflösung von Gleichungen 1. und 2. Grades sowie durch die graphische Darstellung der Funktionen x^n für die einfachsten rationalen Werte des Exponenten Rechnung getragen. Zahlreiche Beispiele dienen zur Einübung des Lehrstoffes. Gelegentlich sind kurze historische Bemerkungen eingestreut und der Anhang enthält eine kurze Geschichte der Arithmetik. An Einzelheiten sei folgendes bemerkt: Der Funktionsbegriff könnte noch mehr mit der Entwicklung des Zahlbegriffes verknüpft werden. Bei der Lehre von den Verhältnissen und Proportionen würde eine frühere Einführung des Proportionalitätsfaktors die Ableitung mancher Sätze vereinfachen. Die auf Seite 210 gegebene Definition ist auch vom methodischen Standpunkte anfechtbar, da sie keineswegs die Schwierigkeiten beseitigt, welche bei der Begründung der irrationalen Zahlen im Unterricht auftreten. In der Lehre von der Wurzel würde es sich empfehlen, die Potenzen mit gebrochenen Exponenten nicht erst am Schlusse, sondern schon früher einzuführen. Die komplexen Zahlen werden zwar auf der Gaussschen Zahlenebene veranschaulicht, es sollte aber auch bei der Entwicklung der Rechengesetze von dieser Veranschaulichung Gebrauch gemacht werden.

Dintzl.

Mues-Schwarz, Mathematisches Unterrichtswerk. Arithmetik und Algebra. Heft 1, Geometrie-Heft 1; Arbeitsbuch für die Mittelklassen höherer Lehranstalten, herausgegeben von W. Escher, B. Reissmann, H. Rüping und W. Schwarz. Aschendorffsche Verlagsbuchhandlung Münster i. W. 1927.

Das Lehrbuch der Arithmetik behandelt die vier Grundoperationen im Gebiete der rationalen Zahlen. An der Hand besonderer Zahlenbeispiele, aus dem dem Schüler schon von früher her geläufigen Rechnen mit besonderen Zahlen werden die allgemeinen Rechengesetze entwickelt und auf das Rechnen mit allgemeinen Zahlen angewendet. Die Auswahl und die Formulierung der aufgestellten Gesetze ist nur allzu einseitig von dem Bestreben beherrscht, die mechanische Rechenfertigkeit zu üben. Methodisch gut durchgeführt ist die enge Verknüpfung der Entwicklung der Rechenoperationen, besonders der inversen Operationen mit der Auflösung linearer Gleichungen sowie die Art, wie auf dieser Stufe der Funktionsbegriff behandelt wird.

Im Lehrbuch der Geometrie werden die wichtigsten Sätze aus der Planimetrie und einige Sätze aus der Stereometrie vornehmlich auf Grund der Anschauung entwickelt. Gelegentlich werden auch Beweise gegeben und besonderes Gewicht wird auf die Klarstellung der Begriffsumfänge bei der Einteilung der oben Figuren gelegt. Weniger didaktischer Wert wird anscheinend den Umkehrungssätzen beigemessen. Die Einteilung des Stoffes sowie die Auswahl der ausgesprochenen Sätze ist eine gute. Im Besonderen sei bemerkt, daß die Parallelenlehre auf dem Satze von der Winkelsumme im Dreieck, der als Grundsatz vorangestellt wird, aufgebaut ist. Es soll dann aber auch der Satz, daß durch einen Punkt außerhalb einer Geraden zu dieser nur eine Gerade gezogen werden kann, als Folgesatz erwähnt werden, da von ihm öfter Gebrauch gemacht wird. Den Forderungen des Arbeitsunterrichtes, welche die Bearbeiter dieses Lehrbuches zu erfüllen suchen, könnte durch eine Vermehrung der Aufgaben besonders aus den Anwendungsgebieten noch mehr Rechnung getragen werden.

Dintzl.

Bauer-v. Hanxleden-Krause, Lehrbuch der Mathematik für Lyzeen und die Unterstufe von Oberlyzeen und Studienanstalten in 7. Auflage bearbeitet von M. Krause. I. Teil: Geometrie. RM 3,—. II. Teil Arithmetik. RM 3,60. Bei Vieweg, Braunschweig 1928.

Der erste Teil umfaßt die Planimetrie, Stereometrie und die ebene Trigonometrie bis zur Berechnung des rechtwinkligen Dreiecks. Er zeichnet sich vor allem durch die äußere Ausstattung des Buches aus (Hervorhebung der Lehrsätze durch Fettdruck, viele instruktive, teilweise zweifarbige Figuren, gutes Papier); man erkennt aber auch überall, daß das Buch aus den Unterrichtserfahrungen von tüchtigen Methodikern hervorgegangen ist. Charakteristisch sind die Vorübungen und Vorbetrachtungen — meist in Form von Konstruktionsaufgaben — welche jedem Lehrsatz vorangestellt sind. Sehr ausführlich und gut durchgearbeitet ist insbesondere der Abschnitt über den Kreis. Wie jedes deutsche Lehrbuch, läßt auch dieses die methodischen Schwierigkeiten im Aufbau der Planimetrie erkennen, der einerseits nicht streng symmetrisch und mit besonderer Beschränkung auf die Anschauung erfolgen, andererseits aber so beschaffen sein soll, daß die Schüler das Beweisen lernen. So steht schon auf Seite 15 der Satz: „Ein Lehrsatz ist ein Satz, dessen Richtigkeit durch Verbindung bereits bekannter Lehrsätze und Grundsätze bewiesen werden muß.“ Dann folgen aber viele Sätze, welche als Lehrsätze ausgesprochen werden und abschließend aus der Anschauung gewonnen werden.

Der zweite Teil behandelt die Operationen der 1., 2. und 3. Stufe, die linearen und quadratischen Gleichungen und ihre graphische Auflösung mit Hilfe der Darstellung der zugehörigen Funktionen. Die Gleichungen sind von Anfang an mit den Rechenoperationen und ihren Gesetzen verbunden. Es wird nur zu früh das mechanische Lösen von Gleichungen eingeführt. Von Wert sind die zahlreichen Aufgaben, besonders die eingekleideten Gleichungen. In der Lehre von den Verhältnissen und Proportionen wird mit Erfolg der Proportionalitätsfaktor zum Beweise der Sätze verwendet. Der Einführung in den Funktionsbegriff an der Hand der graphischen Darstellung empirischer Funktionen sowie der linearen Funktion und ihrer Anwendung auf die graphische Lösung von linearen Gleichungen sind zwei Kapitel gewidmet. Sehr gut ist der im § 64 durchgeführte Gedanke, die graphische Lösung von Gleichungen auf die Lösung von Beweisaufgaben anzuwenden. Die Einführung der Potenzen mit gebrochenen Exponenten erfolgt sehr spät, nämlich am Ende des Wurzelrechnens.

A. Jungbluth, Mathematischer Arbeitsunterricht. P. Henkler. Arbeitsschulmäßiger Rechenunterricht. E. Günther, Physikalischer Arbeitsunterricht. 9. Heft des von A. Jungbluth herausgegebenen Handbuches des Arbeitsunterrichtes für höhere Schulen. Bei Diesterweg, Frankfurt a. M. 1926.

In diesen drei Aufsätzen wird sowohl in allgemeinen Bemerkungen als auch an der Hand besonderer Beispiele gezeigt, wie ein arbeitsschulmäßiger Unterricht in Mathematik und Physik an den höheren Lehranstalten gestaltet werden kann. Alle drei Verfasser stehen erfreulicherweise nicht auf dem Standpunkt der extremen Verfechter des Arbeitsunterrichtes. Es wird u. a. betont, daß die Forderung eines arbeitsschulmäßigen Rechnens schon eine alte sei und es heute nur gelte, diese zu vertiefen. Nicht gerechtfertigt erscheint mir nur die entschiedene Ablehnung des fragend-entwickelnden Verfahrens, von dem zuerst ein Zerrbild entworfen wird, gegen welches dann heftig angekämpft wird. Übertrieben erscheint auch die Wertung der heute so beliebten Art der Formulierung einer Aufgabe, wobei oft nur die gegebenen Größen mitgeteilt werden, ohne daß genau angegeben wird, was zu bemerken sei.

Dintzl.

Verzeichnis von Büchern, die bei der Redaktion eingelaufen sind und einer späteren Besprechung vorbehalten bleiben:

- Berichte der Sächsischen Akademie der Wissenschaften.** 80. Bd. Bei Hirzel, Leipzig 1928. RM 2,25.
- Beyer R.,** Einführung in die Kinematik. Bei Verlagsbuchhandlung M. Jäncke, Leipzig 1928. RM 8,70.
- Brill, Mechanik.** Bei Oldenbourg, München. Geh. RM 18,—, geb. RM 20,—.
- Courant R.,** Differential- und Integralrechnung. Bd. II. Bei Springer, Berlin 1928. Geb. RM 18,60.
- Cramer, Abriß der Finanzmathematik.** Bei Palm & Enke, Erlangen. RM 4,40.
- Fréchet M.,** Les espaces abstraits. Bei Gauthier-Villars, Paris 1928.
- Graff, Grundriß der Astrophysik.** Bei Teubner, Leipzig 1928. Geb. RM 45,—.
- Hartung, Die physikalische Energie.** Bei Dierig, Obersalzenbrunn, Bezirk Breslau, 1928.
- Hauffe G.,** Die symbolische Behandlung der Wechselströme. (Sammlung Göschen 991.) Bei de Gruyter, Berlin 1928. RM 1,50.
- Haupt O.,** Einführung in die Algebra. 2 Bd., (Mathematik in Monographien und Lehrbüchern, Bd. V, 1, 2.) Bei der Akad. Verlagsgesellschaft, Leipzig 1928. Bd. I geh. RM 24,—, geb. RM 26,—; Bd. II geh. RM 20,—, geb. RM 21,50.
- Henkel O.,** Graphische Statik. 2. Aufl. (Sammlung Göschen 695.) Bei de Gruyter, Berlin 1928. RM 1,50.
- Hölder O.,** Die Arithmetik in strenger Begründung. 2. Aufl. Bei Springer, Berlin 1928. RM 3,60.
- Junge, G.,** Einführung in Wesen und Wert der Mathematik. (Sammlung Wissen und Wirken, 56.) Bei Braun, Karlsruhe 1928. RM 3,—.
- Kaluza, Vollmonotone Funktionen.** Bei Niemeyer, Halle a. S. 1928. RM 1,50.
- Kamke E.,** Mengenlehre. (Sammlung Göschen, 999.) Bei de Gruyter, Berlin 1928. RM 1,50.
- Klein F.,** Nichtenklidische Geometrie. Bei Springer, Berlin 1928. RM 18,—, geb. RM 19,50.
- Lecher, Lehrbuch der Physik.** Bei Teubner, Leipzig 1928. RM 18,—.
- Loyarte R. G.,** Fisica General. I. Universidad de La Plata 1927. Preis 8 Dollar.
- Mahler G.,** Physikalische Aufgabensammlung. (Sammlung Göschen, 243.) Bei de Gruyter, Berlin 1927. MR 1,50.
- Michaelis, Einführung in die Mathematik.** 3. Aufl. Bei Springer, Berlin 1928. Geb. RM 18,—.
- Mises R.,** Wahrscheinlichkeit, Statistik und Wahrheit. Bei Springer, Berlin 1928. RM 9,60.
- Neumann, Gesammelte Werke.** Bd. I. Bei Teubner, Leipzig 1928. RM 50,—.
- Orthner, Über physikalische und mathematische Abhängigkeit.** Bei Deuticke, Wien. S 2,40.
- Peters J.,** Sechstellige Tafeln der trigonometrischen Funktionen. Bei F. Dümmler, Berlin 1929. 4^o. VIII + 293 Seiten. RM 48,—, geb. RM 52,—.
- Planck M.,** Einführung in die Theorie der Elektrizität und des Magnetismus. 2. Aufl. Bei Hirzel, Leipzig 1928. Geb. RM 8,—.

- Planck M.**, Einführung in die allgemeine Mechanik. 4. Aufl. Bei Hirzel, Leipzig 1928. Geb. RM 8,—.
- Pringsheim, Fluoreszenz.** (Sammlung Struktur der Materie, Bd. VI.) Bei Springer, Berlin 1928.
- Probleme der modernen Physik.** Sommerfeld-Festschrift. Herausgegeben von P. Debye. Bei Hirzel, Leipzig 1928.
- Richter-Bozen G.**, Die Kraft als fünfte Dimension. Bei Braumüller, Wien 1928. RM 3,50.
- Rose**, Die Schulung des Geistes durch den Mathematik- und Rechenunterricht. Bei Teubner, Leipzig 1928. RM 6,80.
- Salkowski**, Grundzüge der darstellenden Geometrie. (Mathematik in Monographien und Lehrbüchern, Bd. III.) Bei der Akad. Verlagsgesellschaft, Leipzig 1928. Geh. RM 5,75, geb. RM 6,50.
- Schrödinger E.**, Vorlesungen über Wellenmechanik. Bei Springer, Berlin 1928. RM 3,90.
- Sierpinski W.**, Leçons sur les nombres transfinis. Bei Gauthier-Villars, Paris 1928.
- Study E.**, Denken und Darstellung. 2. Aufl. Bei Vieweg, Braunschweig 1928.
- Timerding H. E.**, Zeichnerische Geometrie. (Sammlung Mathematik und ihre Anwendungen, Bd. IV.) Bei der Akad. Verlagsgesellschaft, Leipzig 1928. Geh. RM 25,—, geb. RM 26,—.
- Vivanti G.**, Elemente der Theorie der linearen Integralgleichungen. Übersetzt und mit Anmerkungen versehen von F. Schwank. Bei Helwing, Hannover 1929. Geb. RM 16,60.
- Walther**, Einführung in die mathematische Behandlung naturwissenschaftlicher Fragen. I. Bei Springer, Berlin 1928. Geb. RM 9,60.
- Webster A. E.**, Partial differential equations of mathematical physics. Bei Teubner, Leipzig 1927. Geb. RM 25,—.
- Weyl H.**, Gruppentheorie und Quantenmechanik. Bei Hirzel, Leipzig 1928. RM 22,—.

